

Secondary Publication



Mirbabaie, Milad; Langer, Marie; Rieskamp, Jonas; Hofeditz, Lennart

Transparency Fallacy : Perceived Fairness in Algorithmic Management

Date of secondary publication: 31.07.2026

Version of Record (Published Version), Article

Persistent identifier: urn:nbn:de:bvb:473-irb-116517x

Primary publication

Mirbabaie, Milad; Langer, Marie; Rieskamp, Jonas; Hofeditz, Lennart (2026): Transparency Fallacy : Perceived Fairness in Algorithmic Management, in: Business & information systems engineering : BISE, Wiesbaden: Springer Gabler, Vol. 68, No. 4, pp. 829–850, doi: 10.1007/s12599-025-00963-1.

Legal Notice

This work is protected by copyright and/or the indication of a licence. You are free to use this work in any way permitted by the copyright and/or the licence that applies to your usage. For other uses, you must obtain permission from the rights-holders.

This document is made available under a Creative Commons license.



The license information is available online:

<https://creativecommons.org/licenses/by/4.0/legalcode>



Transparency Fallacy: Perceived Fairness in Algorithmic Management

Milad Mirbabaie · Marie Langer · Jonas Rieskamp · Lennart Hofeditz

Received: 15 February 2024 / Accepted: 8 August 2025 / Published online: 22 September 2025
© The Author(s) 2025

Abstract Algorithmic management (AM) can pose serious challenges for workers on digital labor platforms (DLPs), such as exploitation and a lack of transparency. Prior information systems research has characterized these challenges as unfair practices, particularly in terms of platform functionalities and procedures (e.g., for salaries). This work examined how disclosing information about AM's functionalities and procedures affects its perceived fairness. Building on organizational justice theory (OJT), which distinguishes between distributive, procedural, and informational justice, perceived fairness as a measure of justice was used. The authors conducted an online experiment with 234 DLP workers on a self-developed DLP, on which an algorithm allocated suitable tasks. The workers were assigned to three groups with different types of transparency – distributive, informational, and both distributive and informational – and one control group without enhanced transparency. The results suggest that the

provision of transparency in AM practices and decision-making processes neither affects the perception of informational nor both distributive and informational fairness. Instead, it influences how AM's distributive fairness is perceived on DLPs. Although, according to OJT and research on human–computer interactions, individual differences influence perceived fairness regarding transparency measures, factors such as technology affinity and human–computer trust have not been found to be suitable moderators for their interaction. However, the results revealed that these factors were important predictors of perceived fairness, independent of the level of transparency provided.

Keywords Algorithmic management · Transparency · Fairness · Justice · Digital labor platform

1 Introduction

An increasing number of businesses use algorithms to manage workers' tasks on so-called digital labor platforms (DLPs) (Cameron et al. 2023; Kässi et al. 2021; Rani and Furrer 2021), of which platforms such as Uber, Airbnb, and Amazon Mechanical Turk (Chan and Wang 2018; Möhlmann et al. 2023; Rani and Furrer 2021) are just a few prominent examples. Given their advantages, which include cost-efficient management and flexibility, DLPs are among the fastest-growing markets for labor worldwide (Kässi et al. 2021; Manyika et al. 2016). Due to the opportunities these platforms provide for organizations and workers, the concept of algorithmic management (AM) is increasingly becoming the focus of research (Cram et al. 2022; Zhang et al. 2022). AM refers to “the use of increasingly intelligent algorithms in conjunction with

Accepted after three revisions by Alexander Richter.

M. Mirbabaie (✉) · M. Langer · J. Rieskamp
Faculty Information Systems and Applied Computer Sciences,
University of Bamberg, Kapuzinerstraße 16, 96047 Bamberg,
Germany
e-mail: milad.mirbabaie@uni-bamberg.de

M. Langer
e-mail: marie.langer@uni-bamberg.de

J. Rieskamp
e-mail: jonas.rieskamp@uni-bamberg.de

L. Hofeditz
Hochschule Niederrhein, University of Applied Sciences,
Richard-Wagner-Straße 140, 41065 Mönchengladbach,
Germany
e-mail: Lennart.Hofeditz@hs-niederrhein.de

digital technologies” (Benlian et al. 2022, p. 825) to automate and optimize coordination and control functions traditionally carried out by human managers (Möhlmann 2021; Möhlmann et al. 2021). It is used to perform management responsibilities, such as assigning tasks, communicating with workers, and evaluating workers’ performance (Lee et al. 2015; Schulze et al. 2022), as well as functions with no human involvement on the management side. Parent-Rocheleau and Parker (2022) identified six main AM functions: monitoring, goal setting, performance management, scheduling, compensation, and job termination. The lack of human involvement increases the importance of AM design and configuration for the organizations deploying them (Benlian et al. 2022; Kellogg et al. 2020). Furthermore, previous research has suggested that AM employs two different mechanisms on DLPs: algorithmic matching and algorithmic controlling (Cram et al. 2022; Heinrich et al. 2022; Möhlmann et al. 2021). Möhlmann et al. (2021) defined algorithmic matching as the “coordination of interactions between demand and supply” (p. 2005), which includes using information (e.g., time availability, location, ratings, or evaluation data) to recommend the best matches for both the provider and the platform worker (Heinrich et al. 2022). Algorithmic control refers to the use of “algorithms to monitor platform workers’ behavior and ensure its alignment with the platform organization’s goals” (Möhlmann et al. 2021, p. 2006).

Besides its positive outcomes – such as increased process efficiency and effectiveness (Gal et al. 2017) and flexible work opportunities and autonomy for workers (Heinrich et al. 2022; Wood et al. 2019) – AM can pose serious challenges for organizations and workers, who are the main focus of this research (Benlian et al. 2022; Gal et al. 2020; Parent-Rocheleau and Parker 2022). While algorithms have been implemented to efficiently manage and allocate workforces, issues regarding the treatment of workers and related perceived unfairness have been identified (Köchling and Wehner 2020). Furthermore, disadvantages can arise from the complexity of AM use in organizations, since many parties are involved in the process (Faraj et al. 2018; Heinrich et al. 2022). Thus, AM implementation raises some serious threats, such as power asymmetry (Benlian et al. 2022), information asymmetry (Zhang et al. 2022), and a low level of transparency (Cram et al. 2022; Parent-Rocheleau and Parker 2022) – referred to as opacity (Bujold et al. 2022; Heinrich et al. 2022). Previous research has shown that, compared to prior control systems, AM is “often more opaque in terms of how it directs, evaluates, and disciplines workers” (Kellogg et al. 2020, p. 387). This opacity may stem from a lack of disclosure regarding the design features and source codes of algorithms (Schulze et al. 2022). Emerging concerns about

a lack of transparency and the exploitation of workers can negatively affect perceptions of the fairness of AM (Parent-Rocheleau and Parker 2022; Schulze et al. 2022).

To understand and conceptualize perceived fairness and its constituents, we employ organizational justice theory (OJT) (Colquitt et al. 2001; 2013). Based on the literature, we recognize the varying terminology and interchangeable use of the terms “fairness” and “justice” (Colquitt 2001; Ganegoda et al. 2015; Krasnova et al. 2014; Morse et al. 2022). As Morse et al. (2022) put it, “Although differences exist among the concepts, both are geared toward promoting equity and avoiding bias” (p. 1086). Colquitt et al.’s (2013) seminal work resolved this conflict by defining justice as perceived fairness, which we follow in this study. In the context of AM, a lack of perceived fairness can arise, for example, from a lack of transparency when relevant information is withheld from workers on a platform (Zhang et al. 2022). In fact, a lack of transparency is considered a key characteristic of algorithmic systems, which means that the algorithm’s inputs are unknown, relations between input and output are unknown, and no further explanation is provided (Barredo Arrieta et al. 2020; Burrell 2016; Langer and König 2023; Sokol and Flach 2020). Rewards or payments are often intentionally undisclosed to prevent workers from manipulating systems (Cram et al. 2022; Kellogg et al. 2020; Möhlmann et al. 2021). Theoretically, these issues pertain to the dimensions of *distributive fairness* and *informational fairness*. Distributive fairness refers to perceived fairness in the distribution of rewards or monetary compensation on DLPs, whereas informational fairness relates to the provision of information about AM. Given the opacity of AM practices, studies have suggested that increasing transparency may help mitigate concerns about perceived fairness (Benlian et al. 2022; Cameron et al. 2023; Cheng and Foley 2019; Robert et al. 2020).

Following Diakopoulos and Koliska (2017), we define algorithmic transparency as the “disclosure of information about algorithms to enable monitoring, checking, criticism, or intervention” by workers (p. 811). Previous research has identified that algorithmic transparency influences the perception of fairness in general (Bitzer et al. 2023). However, this paper offers a new perspective by providing a detailed understanding of how different types of algorithmic transparency influence diverse perceptions of fairness. Investigating these effects in the context of AM on DLPs is therefore of particular importance, since perceptions of fairness are highly dependent on context (Jabagi et al. 2024). While transparency can lead to higher perceptions of fairness in some areas, this might not be the case in the context of AM, since DLP structures involve other factors that may also play a role. This aligns with Benlian et al. (2022), who suggested that those results may

not be directly transferrable to AM and that the effects of transparency need to be scrutinized in detail. Furthermore, the AM context requires new forms of transparency, such as the disclosure of evaluation processes or decision-making criteria, when assigning tasks to DLPs (Bujold et al. 2022; Jiang et al. 2023). These forms of transparency may differ in how they affect fairness perception from other types examined in previous algorithmic transparency research. Hence, we pose the following research question:

How does algorithmic transparency affect the perceived fairness of algorithmic management?

We examine the impact of transparency on perceived fairness through an OJT lens. Prior research has highlighted recurring fairness concerns in AM, particularly in distributive and informational dimensions (e.g., Rani and Furrer 2021; Jabagi et al. 2019). Based on this, we focus on corresponding types of transparency and investigate how informational and distributive transparency¹ measures affect workers' perceptions of fairness. The study was implemented using a 4×1 between-subjects design. The participants were divided equally into three experimental groups, each exposed to a different level of transparency – from no transparency to distributive, informational, and a combination of distributive and informational transparency – and a control group with no transparency. Based on this, it was determined whether workers' perceived fairness could be improved with each type of transparency. Finally, the participants were asked about their perceptions of the fairness of the scenarios they had been presented with.

With this research, we contribute to the information systems literature by enhancing our understanding of the use of AM on DLPs, thereby encouraging further research in this field. Our findings enhance our understanding of how transparency influences workers' perceived fairness in digital businesses. Furthermore, this research advances the discourse on algorithmic transparency shaping fairness perceptions (Bitzer et al. 2023) by offering a granular understanding of this relationship and by providing insights into the impact of different types of transparency on the perceived fairness of algorithms. In particular, the impact of each level of transparency on perceived fairness provides new insights into how the relationship between workers and AM on DLPs can be optimized. In this way, we contribute to advancing labor practices by offering guidance to organizations on how to implement transparent information and communication in the control and coordination processes of AM on DLPs. Our study offers insights into workers' needs by determining the extent to which fairness is perceived on DLPs and how transparency can help improve this perception. Practitioners will

understand the extent to which explanation-based transparency improves workers' perceived fairness on DLPs that utilize AM and how this can be leveraged for the benefit of their organizations and workers. Practitioners that use DLPs can use our study's findings to enhance the relationship between workers and the algorithms they utilize for AM. In this way, this research will help extend the knowledge of the possible applications of AM in organizational contexts.

2 Theoretical Background and Hypotheses

2.1 The Link Between Fairness and Transparency in Algorithmic Management

When describing organizational fairness, research relies on OJT (Colquitt et al. 2001; Cui et al. 2023), which aims to explain how workers perceive fairness in the workplace (Greenberg 1990). To build on this, research on OJT has distinguished between various dimensions of fairness, such as distributive, procedural, and interactional fairness (Colquitt 2001). While overall fairness refers to the perception and evaluation of a system as a whole (Beugre and Baron 2001), distributive, procedural, and interactional fairness refer to the different dimensions that constitute fairness. Since workers on DLPs desire systems that are fair in terms of outcomes and information handling (Beugre and Baron 2001; Zhou et al. 2023), in this study, we consider distributive and interactional fairness. The distinction between these dimensions of fairness is particularly important in AM, as unfairness can arise at various touchpoints, such as nontransparent communication between the algorithm and workers or the unfair distribution of rewards by the algorithm (Rani and Furrer 2021; Zhou et al. 2023). Although the three dimensions of fairness are linked to positive outcomes, they differ conceptually (Cui et al. 2023). Examining each dimension in depth can, therefore, have a major impact on workers' perceptions of fairness and lead to different outcomes (Colquitt 2001; Greenberg 1990). Each individual dimension has a distinct impact on factors such as system trustworthiness and privacy concerns (Krasnova et al. 2014). Therefore, in this study, we distinguish between distributive and interactional fairness to analyze how they are affected by transparency.

Distributive fairness refers to the extent to which equal contributions (e.g., work time) yield equal outcomes (e.g., points) (Alexander and Ruderman 1987; Krasnova et al. 2014; Robert et al. 2020). Therefore, it can be defined as the extent to which the distribution of an outcome is perceived as fair (Lu et al. 2024). In the context of DLPs, workers seek fair rewards for their performance, which is

¹ Creating transparency by providing information on and explanations about systems, processes, or decisions (Felzmann et al. 2019).

determined by comparing the time worked to the points awarded. Moreover, the relationship between workload and reward is particularly important in assessing the fairness of DLPs. Therefore, we capture the importance of distributive justice in the context of DLPs from the importance that workers attach to their efforts being rewarded by a fair assessment of points (Zhou et al. 2023). To enable distributive justice, the concept of transparency can be utilized. Accordingly, the disclosure of rewards allocated by AM related to working time can be attributed to distributive transparency, which can thus be defined as openness in the process of distributing rewards (e.g., point allocation) (Krämer and Wiewiorra 2015; Sikayu et al. 2022). Previous research has generally linked (distributive) transparency positively to distributive fairness; thus, the provision of transparency has an impact on distributive fairness (Bujold et al. 2022). This may be because transparency shows how rewards are distributed in relation to individual contributions, thereby enabling workers to assess the fairness of such distribution (Bujold et al. 2022). Therefore, we propose the following hypothesis:

H1a: Providing distributive transparency increases the perceived distributive fairness of AM.

Interactional fairness focuses on organizations' treatment of their workers (Luo 2007) and can be separated into interpersonal fairness (i.e., quality of interpersonal communication between different stakeholders) and informational fairness (i.e., explanations about AM processes and outcomes) (Colquitt 2001; Robert et al. 2020). Since interpersonal communication does not apply to AM, we focus only on the aspect of fair interactions that involves keeping workers well informed (i.e., informational fairness) (Cohen-Charash and Spector 2001; Krasnova et al. 2014; Robert et al. 2020). Therefore, informational fairness can be defined as accurately explaining information management practices in AM (Colquitt 2001). In our context, workers desire timely information on AM processes and technical functions and how these affect the allocation of work on DLPs. Therefore, we capture the importance of informational justice in the context of DLPs through the importance that workers attach to AM processes and technical functions that are reliable, understandable, and visible. If a DLP gives its workers access to information on AM decision-making processes and functions, this can be described as establishing transparency (Lu et al. 2014). We refer to informing workers about how data are used as informational transparency, which is defined as "a means to make users more informed" about the information management processes of AM on DLPs (Jiang et al. 2023, p. 1695). Consequently, disclosing information on a DLP that addresses the algorithm's processes, outcomes, and information handling is expected to increase perceived

informational fairness (Colquitt 2001; Krasnova et al. 2014). Therefore, we propose the following hypothesis:

H1b: Providing informational transparency increases the perceived informational fairness of AM.

In general, AM tends to contribute to ethical challenges (Gal et al. 2020) because of its opaque and incomprehensible design (Kellogg et al. 2020; Zhang et al. 2022). Because of this, questions have been raised about how AM affects workers on DLPs. Algorithmic unfairness and injustice² are two major outcomes that DLP workers face due to AM (Bujold et al. 2022; Geissinger et al. 2022; Schulze et al. 2022). Hence, we define algorithmic unfairness as "AM practices that give a rise to systematic disadvantage for [platform] workers" that results from algorithmic decision-making (Schulze et al. 2022, p. 2). To decrease the negative effects of AM, studies have suggested various approaches, with algorithmic transparency (Langer and König 2023; Parent-Rocheleau and Parker 2022; Rani and Furrer 2021; Schulze et al. 2022; Zhang et al. 2022) and the perception of fairness (Bujold et al. 2022; Kordzadeh and Ghasemaghaei 2022; Krasnova et al. 2014; Marjanovic et al. 2022) being the most common. In AM, fairness mainly concerns workers' perceptions of interactions with it and the decisions it makes (Lee 2018). These perceptions are largely shaped by whether users understand how an algorithm makes decisions. Understanding this process is difficult due to its underlying complexity (Cameron et al. 2023) and the large amount of data used for decision-making (Faraj et al. 2018; Tarafdar et al. 2022). This can cause workers to perceive systems based on AM as unfair (Parent-Rocheleau and Parker 2022). Thus, perceived fairness requires the disclosure of information about the functions and procedures or performance measurement logic of AM (Schulze et al. 2023). In conclusion, AM causes friction for workers, often because of a lack of transparency of DLPs (Gal et al. 2020; Kellogg et al. 2020), which often do not disclose the necessary information "to enable understanding, critical review, and adjustment" (Bitzer et al. 2023, p. 293) regarding algorithms' decision-making processes (Cameron et al. 2023; Möhlmann et al. 2023).

Previous research has identified information disclosure as a solid means to increase the perceived fairness of AM (Schulze et al. 2023; Spiekermann et al. 2022; Rani and Furrer 2021), but empirical knowledge on the relationship between AM transparency and perceived fairness remains limited. We aim to expand the understanding of how algorithmic transparency can influence the perceived fairness of AM by investigating the extent to which the

² Following Morse et al. (2022), the terms "fairness" and "justice" are used interchangeably in this study.

provision of transparency – that is, providing information about AM processes or the distribution of rewards – leads to higher perceived fairness than when there is no algorithmic transparency. Therefore, we propose the following hypothesis:

H1c: *Providing distributive and informational transparency increases the perceived fairness of AM.*

2.2 The Role of Trust in and Affinity for Technology

OJT suggests that individual differences influence how fairness-related information is processed and evaluated (Colquitt et al. 2013). In the context of DLPs, where AM systems allocate work and structure interactions, individual factors become particularly relevant. Affinity for technology – defined as a personal disposition toward engaging with and understanding technology – may influence how workers interpret transparency measures in algorithmic decision-making (Heilala et al. 2023). Mosaferchi et al. (2023) showed that individuals with high technology affinity tend to place more trust in autonomous systems, particularly when system processes are made transparent. Applied to AM, this suggests that tech-savvy workers may respond more positively to transparency initiatives, perceiving algorithmic processes as fairer and more comprehensible.

Moreover, those with a high affinity for technology may be more capable of understanding and evaluating technical and transparent information. This could influence how they react to the provision of information, leading them to perceive transparency as beneficial to fairness (Jarrahi et al. 2021). Therefore, we argue that affinity for technology moderates the effect between transparency and the perception of fairness and propose the following hypotheses:

H2a: Affinity for technology moderates the effect of distributive transparency on the perception of distributive fairness of AM.

H2b: Affinity for technology moderates the effect of informational transparency on the perception of informational fairness of AM.

H2c: Affinity for technology moderates the effect of distributive and informational transparency on the perception of fairness of AM.

Since most research has focused on how workers on DLPs react to AM, perceptions such as trustworthiness have been the focus, in addition to perceived fairness (Heinrich et al. 2022). Trust is generally an important prerequisite for AM acceptance and use (Moussawi et al. 2021; Schmidt et al. 2020; Zhang et al. 2021). When working with an algorithm in the context of AM, it is

highly important to build trust, which can be defined as the “willingness to be vulnerable to another party based on the belief that the latter party is (1) competent, (2) open, (3) concerned, and (4) reliable” (Mishra 1996, p. 5). This is particularly necessary since it can have an impact on the collaboration between DLP workers and technology (Jarrahi et al. 2021). As with traditional working arrangements, AM relies heavily on worker–manager collaboration and therefore requires a trustful relationship. Here, workers’ trust in the algorithm’s abilities and decisions plays a major role in the success of AM (Benlian et al. 2022). Beyond interpersonal trust, trust in technology itself is a crucial factor, as it influences perceptions of reliability, functionality, and helpfulness (McKnight et al. 2011). According to the integrative trust model (Mayer et al. 1995), trust is shaped by ability, benevolence, and integrity, which also apply to technology-mediated work settings. Trust in an algorithm can therefore shape fairness perceptions by reducing uncertainty and mitigating concerns about transparency (McKnight et al. 2002). Additionally, preexisting trust in AM can have an impact on how workers perceive an algorithm to be fair (Krasnova et al. 2014). Thus, if workers have a high level of trust in technology – such as AM – this could have an impact on their perception of the fairness of the algorithm used on the DLP. Therefore, we argue that trust in AM moderates the effect between transparency and the perception of fairness, from which we derive the following hypotheses:

H3a: Trust in AM moderates the effect of distributive transparency on the perception of distributive fairness of AM.

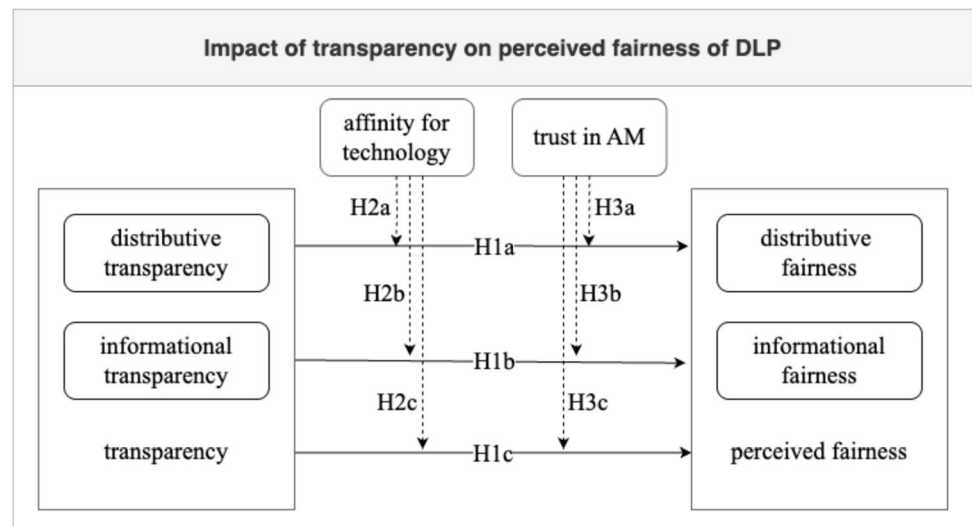
H3b: Trust in AM moderates the effect of informational transparency on the perception of informational fairness of AM.

H3c: Trust in AM moderates the effect of distributive and informational transparency on the perception of fairness of AM.

Figure 1 visualizes our research model.

3 Research Design

To examine the different levels of transparency, we implemented a 4×1 between-subject design. To examine user behavior on a DLP, we conducted an online experiment on a self-developed DLP prototype on which users had to perform typical tasks delegated by simulated AM. We collected quantitative data using survey questions and statements before and after interactions on the prototypical platform and matched these data with those from the users’ interactions. The platform prototype was built using the no-

Fig. 1 Research model

code tool Bubble.io,³ which allows for the simulation of different types of interactive online platforms and tools. In addition, we collected qualitative data by asking for open feedback based on users' perceptions following their interactions on the platform. The questionnaires are available in the online appendix.

As with any study, a key decision involved whether to recruit participants from a student sample or to use a recruitment platform, such as Prolific.⁴ Both options have their own advantages and limitations and reflect different philosophical perspectives on generalizability and contextual relevance. We used Prolific to consider a suitable target group, which allowed us to collect data on a single day (September 9, 2024).

Prolific was an ideal platform for our study because it aligned closely with the characteristics of the DLP we developed and investigated. According to Dunn et al. (2023), Prolific qualifies as a DLP, since it manages task allocation, worker interactions, and remuneration through algorithmic processes. These features make Prolific not only a recruitment tool but also a representative context for studying AM and fairness perceptions in DLP environments.

Furthermore, Prolific offers advanced prescreening options and transparent compensation structures, ensuring high data quality and relevance to our target group of DLP workers. Compared to other platforms – such as MTurk – Prolific demonstrates superior performance in data reliability, as participants consistently pass attention checks and engage more thoughtfully (Douglas et al. 2023). This makes it particularly well suited for experiments requiring

precise group allocation and interaction tracking, such as our 4×1 between-subject design.

By using Prolific, we not only ensured access to a relevant and high-quality sample but also leveraged a platform that mirrored the core dynamics of our study, thereby enhancing the ecological validity of our findings. The participants received appropriate monetary compensation for their time, in line with Prolific's recommendations. Prolific offers a compensation range with different categories, from "low" to "great." For our study, we chose the "fair" category – a mid-level option – to ensure realistic monetary compensation typical of a DLP. We also applied the following prescreening for our participants: they should be employed full- or part-time or work under another type of employment status (e.g., as freelancers).

The participants were distributed equally into four experimental groups, which differed according to the types of transparency displayed to the participants on the recruitment platform used in the study. The first group was not offered any transparency, the second group was provided with distributive transparency, the third group received informational transparency, and the fourth group was exposed to both distributive and informational transparency.

3.1 Manipulation Check

To verify the effectiveness of our treatment conditions, we conducted a prestudy, which served as a manipulation check for our main experiment. To this end, the manipulation check included 120 respondents acquired via Prolific. Here, the four conditions (no transparency, distributive transparency, informational transparency, and both distributive and informational transparency) were displayed to the participants, who were asked to evaluate

³ <https://bubble.io/>.

⁴ <https://www.prolific.co/>.

Table 1 Post Hoc Mann–Whitney tests

Mann–Whitney test between different groups		U	<i>p</i> -value
No transparency (M = 4.2, SD = 1.7)	Distributive transparency (M = 6.0, SD = 0.9)	2515.0	< 0.000*
No transparency (M = 4.2, SD = 1.7)	Informational transparency (M = 5.0, SD = 1.4)	5214.0	< 0.000*
No transparency (M = 4.2, SD = 1.7)	Distributive and informational transparency (M = 5.9, SD = 1.1)	2810.0	< 0.000*

*Significance at the 0.05 level

their perceptions of AM transparency on the DLP (i.e., the study recruitment platform). Specifically, they were asked to evaluate all four groups' platform designs in terms of their transparency. At the beginning of the prestudy, we presented the different platform designs (i.e., different in terms of transparency) on the same page to create an overview and provided further instructions on how to proceed. In the next step, we displayed a screenshot of each design (i.e., transparency level), and to enable the measurement, we asked the participants to answer Höddinghaus et al. (2021) construct on perceived transparency, which was measured on a seven-point Likert scale (see Table 2). Since one respondent failed the attention checks, we removed their data from the final dataset (N = 119).

The results demonstrated that the four conditions were perceived differently in terms of transparency. To analyze the differences in the perceptions of transparency among the four groups, we first collected descriptive statistics to gain an overview of the tendencies of each group. The data showed a difference between the control group (i.e., no transparency), which had a significantly lower perception of transparency (M = 4.2), and the three treatment groups (distributive transparency: M = 6.0, informational transparency: M = 5.0, and both distributive and informational transparency: M = 5.9). Moreover, the results indicated that the provision of both distributive and informational transparency tended to be perceived as the most transparent. To evaluate whether the differences between the groups were statistically significant, a Kruskal–Wallis test was conducted ($F = 106.84$, $p = 0.000 < 0.05$). Since the *p*-value was below the conventional significance level of 0.05, this suggested that the differences between the groups were statistically significant. Because the Kruskal–Wallis test indicates only whether significant differences exist – but not between which groups – post hoc tests were performed. Therefore, the Mann–Whitney U test was used to compare the groups pairwise. The results showed significant differences between the control group (no transparency) and the three treatment groups (distributive, informational, and both distributive and informational transparency) (see Table 1).

3.2 Materials and Procedure

To examine differences in the perceived fairness of the DLP design, which was enabled by AM, it was important to use a controllable and customizable setting instead of a real-world setting (Hofeditz et al. 2022). Using Prolific, we could ensure the selection of a suitable target group – as its users work in the field of digital labor – and apply different types of prescreening. Previous research has suggested that the Wizard of Oz approach can ensure controllable and reproducible study outcomes, whereas applying a real algorithm could cause too much complexity and variance (Schoonderwoerd et al. 2022; Weiss et al. 2009; Wilson and Rosenberg 1988).

After a briefing, the participants were asked to generate a personal mother code for pseudonymized identification and to provide information on their demographics, such as age, gender, occupation, education, industry, and frequency of working on DLPs, in an online survey using LimeSurvey.⁵ Figure 2 visualizes the overall procedure of the study, which we explain in detail below.

Once the study began and demographic data were entered, the participants were redirected to the study recruitment platform we designed. This platform was designed in such a way that we could apply the experimental conditions to each group. The main task of the participants was to take part in up to three different short dummy studies, which were displayed on our newly developed platform for algorithm-supported study recruitment. All groups were presented with the same dummy studies to ensure consistency. On the platform, they were asked to mention their individual mother codes and then to begin interactions by requesting studies. Figure 3 shows a screenshot of this start screen.

After requesting studies by pressing the “Start” button (see Fig. 3), all groups saw a loading screen showing an (allegedly) algorithm-based calculation of the currently available and suitable studies in the form of a loading bar to simulate the AM allocation process. They were told that the studies on the platform were assigned by an algorithm. Figure 4 shows examples of the loading screen.

⁵ <https://www.limesurvey.org/en/>.

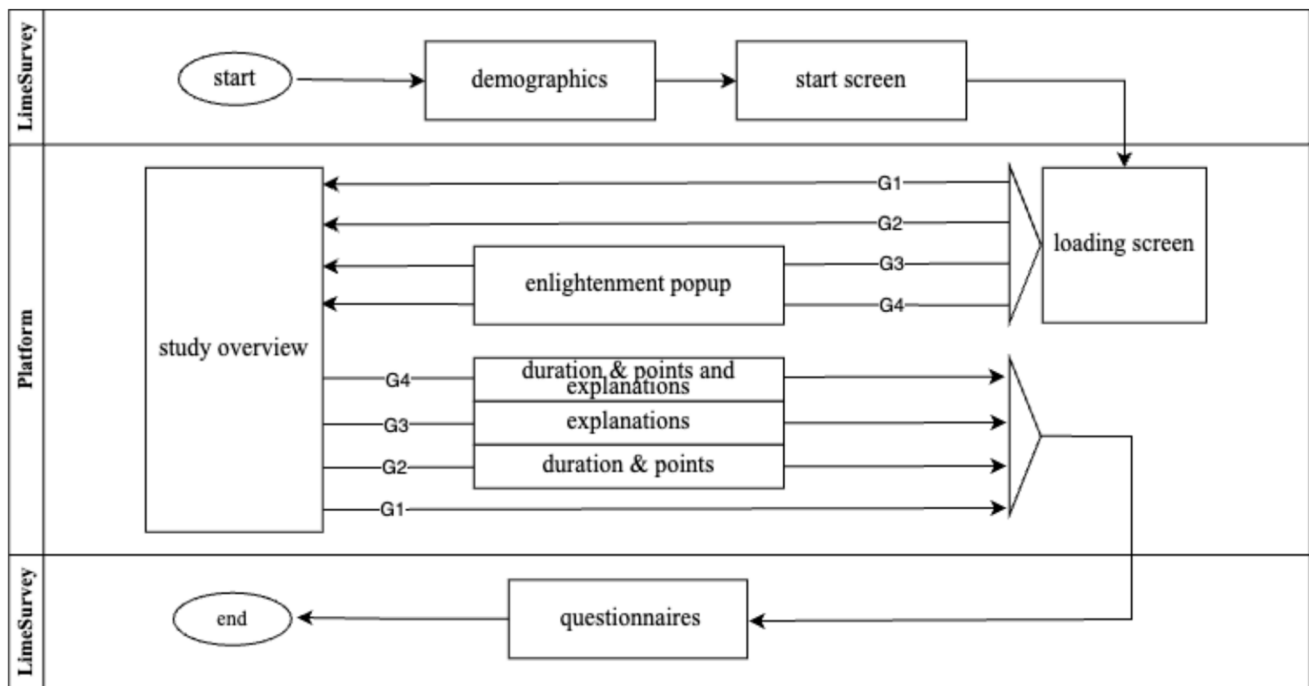


Fig. 2 Study procedure

Fig. 3 Prototype platform start screen

When the loading screen indicated completion, the process varied between the groups. Whereas Groups 1 and 2 were forwarded directly to an overview of the studies that had been identified and assigned by the algorithm, Groups

3 and 4 were first presented with information regarding the algorithmic decision-making process (i.e., informational transparency; see Fig. 2) via a pop-up, which appeared before the participants were redirected to a study overview page. In the pop-up, the participants were given information on how the algorithm operated – that is, its algorithmic matching and control practices. To ensure that the participants did not overlook the information and simply move on without reading it, we included a checkbox that had to be clicked to indicate they acknowledged the information presented (see Fig. 5).

The participants in Groups 3 and 4 were then presented with an overview of the study. On this page, each participant could participate in up to three short surveys. After finishing them, they had to enter a completion code. At any time, they could press the “Checkout” button to stop completing tasks and proceed with our questionnaires. Here, Group 1 was shown only the available studies, without any transparency dimension (see Fig. 6). For Group 2, distributive transparency was added by showing participants the points to be achieved (rewards) and the corresponding duration (contribution), in addition to the studies allocated by AM in the table (see Fig. 7). When the completion code was successfully entered, the points earned were immediately calculated and displayed, helping the participants recognize the points (i.e., their rewards). The participants in Group 3 did not receive any information regarding distribution but instead could retrieve previously received information from the pop-up (i.e., informational

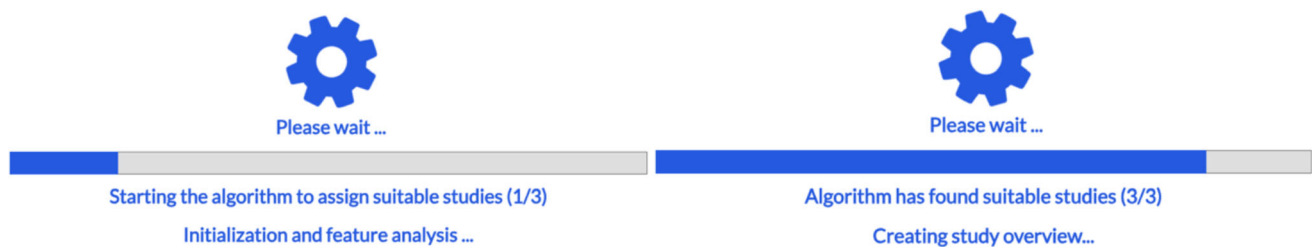


Fig. 4 Exemplary loading screen simulating an AM process for task allocation

Fig. 5 AM informational transparency in a pop-up window (for Groups 3 and 4)

Enlightenment: Algorithmic Management Practices

Assigned Studies

You have been assigned studies by a prototype of an algorithm. This platform uses the algorithm to match you with suitable studies by analyzing your demographic data.

The algorithm works in two ways:

1. **Algorithmic Matching:** Assigns studies based on your data (e.g., age, gender, occupation, participation points) and study availability.
2. **Algorithmic Controlling:** Evaluates feedback from study owners to improve future matches. Participation is optional, and you can stop or reject any study without any disadvantages.

For questions or concerns, please contact us at: [REDACTED]

☐ I understand these algorithmic management practices

Submit

transparency). Therefore, we added an information box to the overview page so that all relevant information regarding the algorithm's allocation procedure could be accessed at any time (see Fig. 8). Group 4 was provided with both distributive and informational transparency on the overview page (i.e., they were shown the duration and points in a table) as well as the option to retrieve the information from the previously displayed pop-up via the information box (see Fig. 9). When the participants pressed the "Checkout" button, they were asked to close the website for the study recruitment platform and go back to the online survey on LimeSurvey (see Fig. 2).

After interacting with the self-developed platform, the participants were asked to complete the Affinity for Technology Interaction (ATI) scale, since technology affinity is a known factor influencing the perception of algorithms (H4a–H4c) (Franke et al. 2017). In addition, the Human–Computer Trust Scale (HCTS; Gulati et al. 2019) was used as a scale for measuring trust in AM (H5a–c).

Next, we presented the items on perceived fairness according to Höddinghaus et al. (2021), which were suitable for our research in the context of AM. As fairness is a multifaceted concept, we also inquired about its distributive and informational dimensions, according to Colquitt (2001). Table 2 provides an overview of the scales.

To ensure that the participants remained focused on the study, we included two attention checks, asking them to select a particular option on a scale. The first attention check required the participants to recall and enter the word "Sun," which was displayed at the end of their interactions on the digital labor platform. The second attention check was embedded within the ATI scale, instructing participants to select "strongly disagree" for a specific item. Those who failed to pass the two attention checks were excluded from the final sample. After the participants completed the scales, they were presented with a debriefing page and redirected to Prolific to receive their monetary compensation. The participants required, on average,

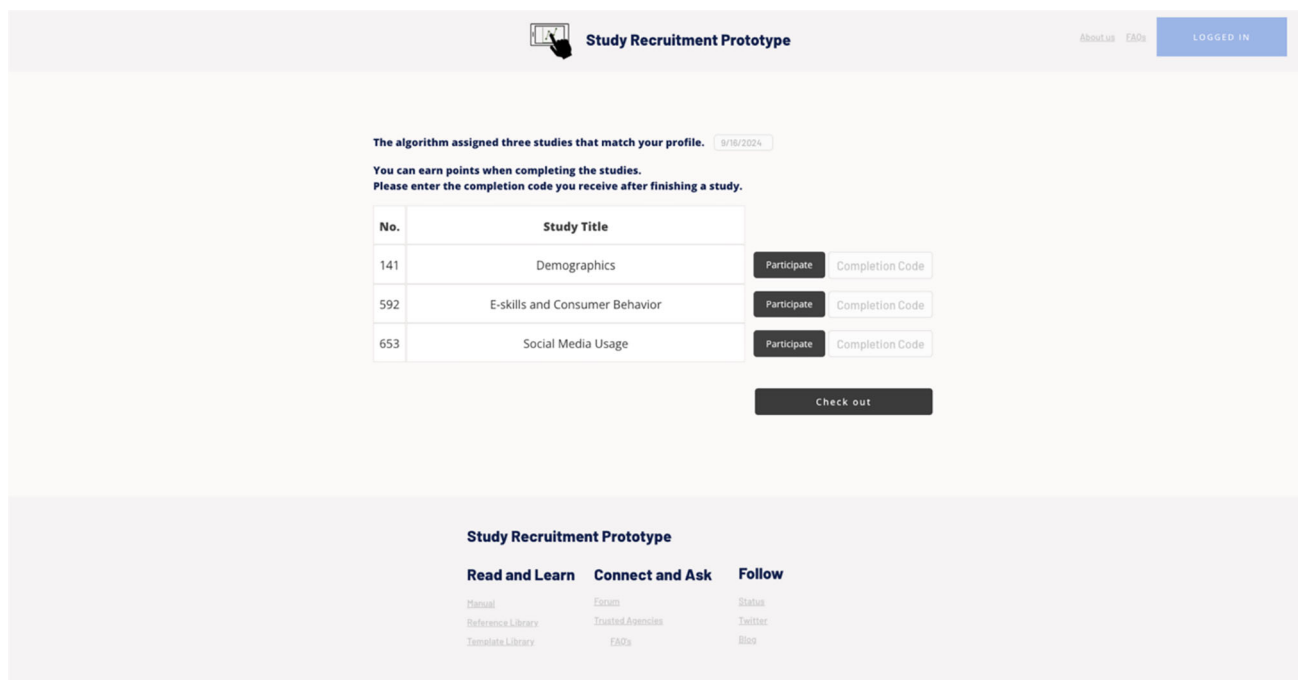


Fig. 6 Example of the study overview (for Group 1)

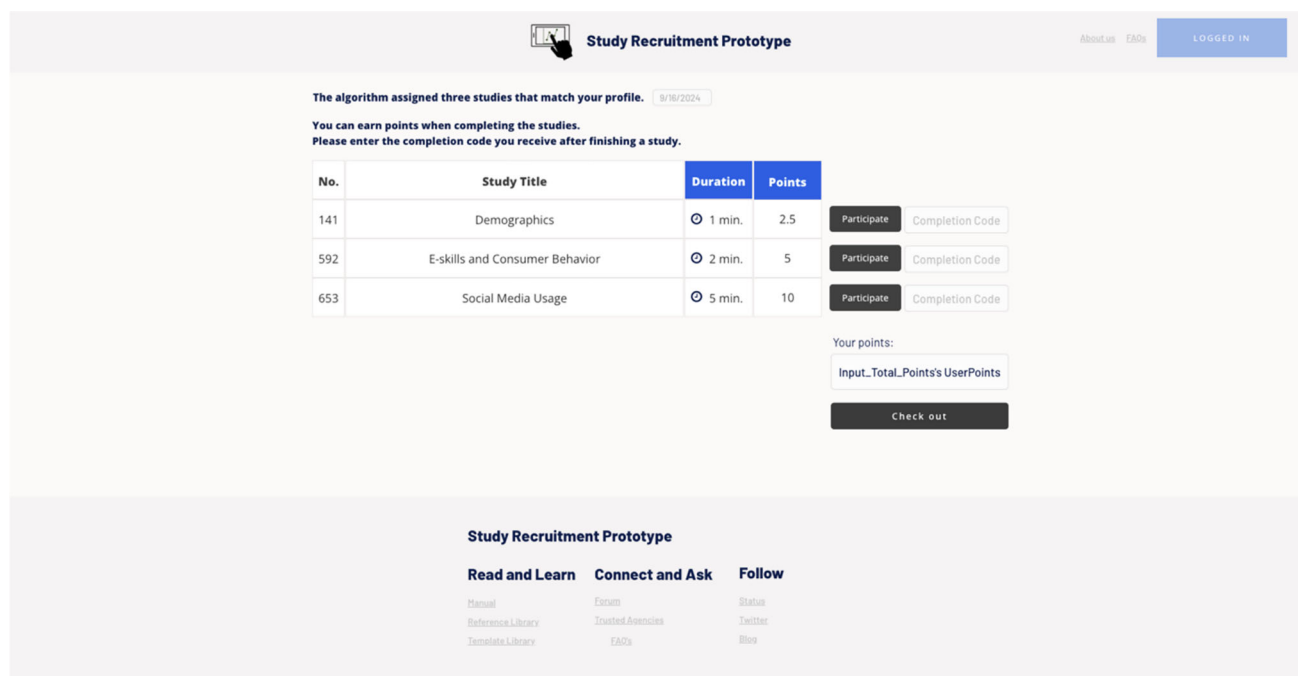


Fig. 7 Example of the study overview, displaying allocated studies, the required time, and the respective points (for Group 2)

21:44 min to complete the study (22:12, 20:00, 23:16, and 21:26 min for Groups 1–4, respectively). Those who completed the study substantially faster than the average also failed the attention checks, leading to their exclusion from our analysis.

3.3 Demographics

A total of 236 participants took part in this study. However, two participants were excluded because they did not pass the attention checks; therefore, 234 valid cases were included in the analysis. To determine the required sample

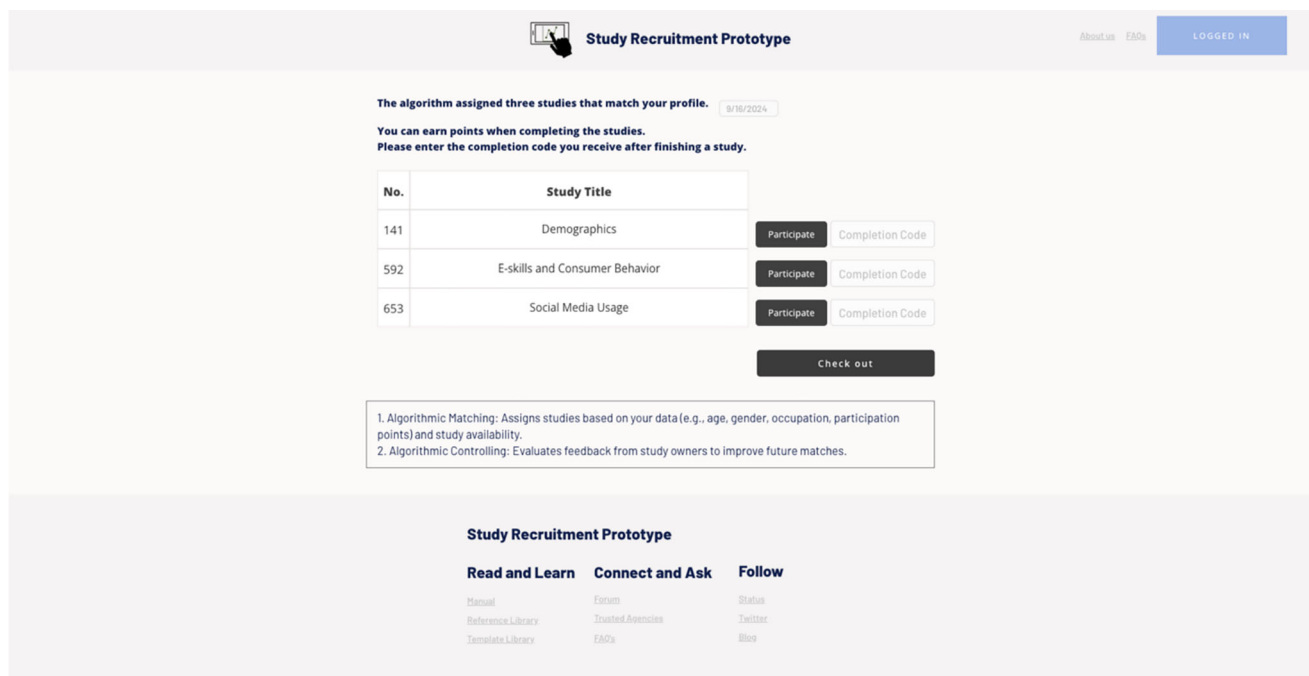


Fig. 8 Example of the study overview, displaying previously received information on the pop-up once again (for Group 3)

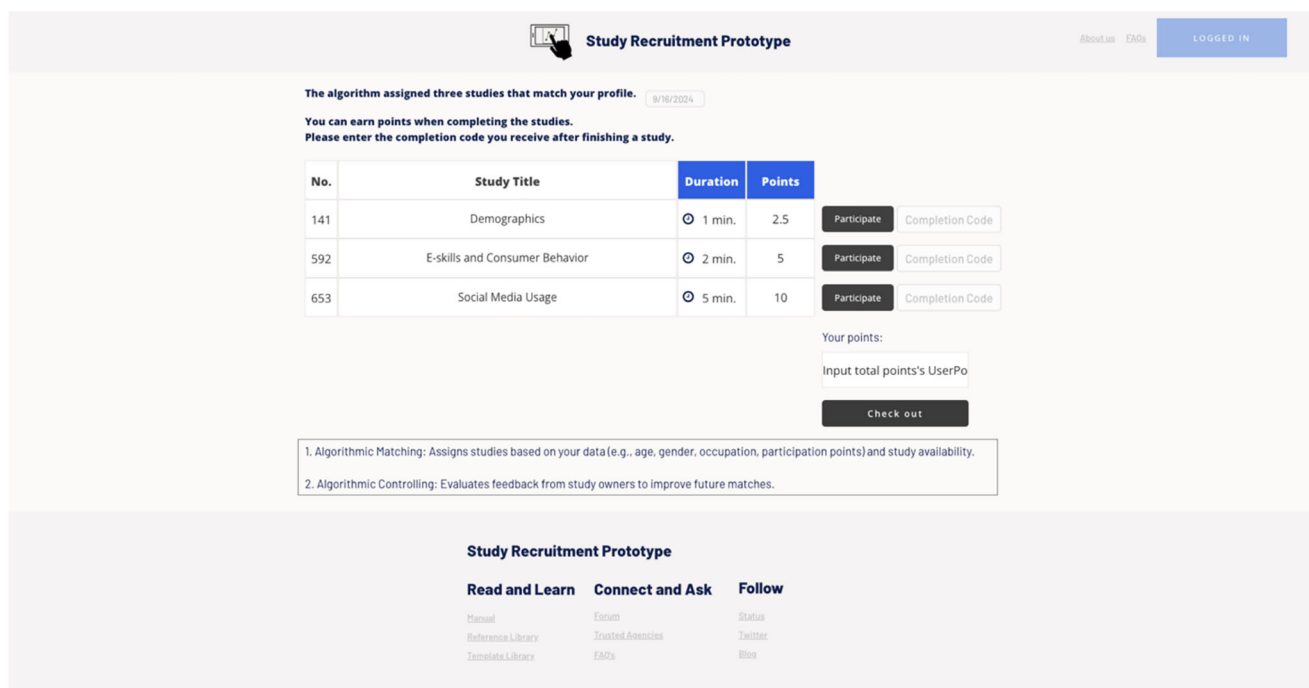


Fig. 9 Example of the study overview, displaying allocated studies, the required time, and the respective points and previously received information from the pop-up (for Group 4)

size, we conducted a G*Power analysis⁶ for the four groups and a medium effect size following standard power analysis procedures, which indicated a required total sample

⁶ <https://www.psychologie.hhu.de/arbeitsgruppen/allgemeine-psychologie-und-arbeitspsychologie/gpower>.

size of 280 participants. However, after data cleaning and applying our exclusion criteria, we retained 234 valid responses, which was slightly below the initially estimated sample size. While this deviation marginally reduced the statistical power, our sample remained within an

Table 2 Questionnaires

Questionnaire	α	Scale	Sources
Affinity for technology interaction	0.90	6-point Likert scale	Franke et al. (2017)
Human–computer trust scale	0.83–0.88	7-point Likert scale	Gulati et al. (2019)
Perceived fairness	0.96	7-point Likert scale	Höddinghaus et al. (2021)
Perceived transparency	0.82	7-point Likert scale	Höddinghaus et al. (2021)
Distributive fairness	0.93	5-point Likert scale	Colquitt (2001)
Informational fairness	0.90	5-point Likert scale	Colquitt (2001)

acceptable range for detecting medium-sized effects. The final sample included 60 participants from Group 1, 56 participants from Group 2, 58 participants from Group 3, and 60 participants from Group 4. They ranged in age from 18 to 72 years ($M = 30.8$, $SD = 9.70$), and 109 were women (46.58%), 121 were men (51.71%), and four were diverse (1.71%). Moreover, the participants disclosed a high level of education. In total, 75.65% had a university degree, while 24.35% had an intermediate leaving degree, an advanced technical college degree, or a higher education entrance qualification. Overall, 82.91% reported being employed, with 66.24% being in full-time employment and 16.67% being in part-time employment. Since we recruited the participants through the Prolific DLP, we were interested in examining how frequently they worked as digital workers. Of the participants, 52.99% indicated that they worked as digital workers daily, while 28.63% reported doing so on a weekly basis, 11.11% monthly, and 4.70% annually.

4 Findings

4.1 Quantitative Findings

4.1.1 Effect of Transparency on Perceived Fairness

To obtain a general overview of the different distributions within the four groups and to further analyze the effects of transparency on perceived fairness, we first examined the descriptive statistics.⁷ The results revealed that, on average, perceived fairness was generally rated higher than the detailed levels of fairness (i.e., information and distributive fairness) across all four groups. Additionally, the average values for distributive fairness were higher across all groups than those for informational fairness; therefore, distributive fairness was perceived by the participants as stronger than informational fairness. Table 3 shows the means and standard deviations of the four conditions and tested scales.

⁷ The data and all analyses are available online at https://osf.io/gwp2q/files/osfstorage?view_only=07b1d8df7af34596b78ab643f3ef2698.

Before analyzing whether providing algorithmic transparency on DLPs leads to a higher perception of fairness of AM (H1a–H1c), we tested our data for normal distribution using the Shapiro–Wilk test. Since the results showed that the assumptions for a normal distribution were not met (i.e., some of the data were not normally distributed; see Table 4), we chose the Kruskal–Wallis test to analyze the differences among the four independent groups (H1a–H1c).

To test for different effects of distributive and informational transparency, we conducted Kruskal–Wallis tests for both distributive and informational fairness (H1a–H1b). First, we analyzed whether providing **distributive transparency** on a DLP leads to higher perception of distributive fairness of AM (H1a). Our results revealed that there were significant differences regarding distributive fairness between the four independent groups ($F = 13.708$, $p = 0.003 < 0.05$). Since the Kruskal–Wallis test provides only insights into whether at least one group differs significantly from the others and does not identify which specific groups differ, we conducted a Mann–Whitney U test to test the specific differences between those groups. Our results showed that there were significant differences between Group 1 (i.e., control group with no transparency) ($M = 3.48$) and Group 2 (i.e., distributive transparency) ($M = 3.85$, $F = 1283.5$, $p = 0.027 < 0.05$), between Group 2 (i.e., distributive transparency) ($M = 3.85$) and Group 3 (i.e., informational transparency) ($M = 3.44$, $F = 2095.5$, $p = 0.007 < 0.05$), and between Group 1 (i.e., control group with no transparency) ($M = 3.48$) and Group 4 (i.e., both distributive and informational transparency) ($M = 3.90$, $F = 1334.5$, $p = 0.014 < 0.05$). To further determine the mean differences between the groups, we calculated the effect sizes according to Cohen's d (Cohen 1988). Therefore, our results showed that the effect size between Groups 1 and 2 (i.e., control group with no transparency and the distributive transparency group, respectively) ($d = -0.39$) and between Groups 1 and 4 (i.e., control group with no transparency and the distributive and informational transparency group) ($d = -0.45$) indicated a small to moderate effect. This means that, on average, the control group (i.e., no transparency) had a lower perception of distributive fairness than the groups

Table 3 Descriptives results

No	Condition	Perceived fairness		Distributive fairness		Informational fairness	
		M	SD	M	SD	M	SD
1	No transparency	4.76	1.38	3.48	0.96	3.17	1.09
2	Distributive transparency	4.86	1.09	3.85	0.93	3.33	1.11
3	Informational transparency	4.87	1.11	3.44	0.85	3.19	1.00
4	Distributive and informational transparency	4.95	1.32	3.90	0.87	3.61	0.95

Perceived fairness was measured using a seven-point Likert scale, while Distributive Fairness and Informational Fairness were measured using a five-point Likert scale

Table 4 Shapiro–Wilk test for normal distribution

No	Condition	Perceived fairness		Distributive fairness		Informational fairness	
		F	<i>p</i> -value	F	<i>p</i> -value	F	<i>p</i> -value
1	No transparency	0.96	0.06	0.96	0.04*	0.97	0.13
2	Distributive transparency	0.94	0.01*	0.93	0.003*	0.94	0.005*
3	Informational transparency	0.95	0.03*	0.96	0.06	0.96	0.04*
4	Distributive and informational transparency	0.95	0.01*	0.94	0.004*	0.96	0.06

*Significance at the 0.05 level

that received distributive transparency. These results underline the effect of the provision of distributive transparency on the participants' perceptions of distributive fairness, since this condition differed significantly from the other transparency levels (i.e., no transparency and informational transparency). Thus, it can be concluded that the provision of distributive transparency has a significant – although only small to moderate – effect on the perception of distributive fairness compared to the other transparency levels. Moreover, our results showed no significant differences between Group 2 (i.e., distributive transparency) ($M = 3.85$) and Group 4 (i.e., both distributive and informational transparency) ($M = 3.9$) regarding distributive fairness. These results highlight that the provision of distributive transparency in both groups (i.e., Groups 2 and 4) did not lead to different results, since we were able to confirm the effect in both groups. Therefore, the effect of distributive transparency on the perception of perceived distributive fairness occurred equally in both groups and did not differ significantly ($F = 1650.5$, $p = 0.872 > 0.05$). However, we found statistically significant differences between providing no transparency and providing distributive transparency in the perception of distributive fairness; therefore, we found support for H1a. Second, we analyzed whether providing **informational transparency** on a DLP leads to a higher perception of informational fairness of AM (H1b). Our results showed that there were

no significant differences regarding informational fairness between the four independent groups ($F = 6.880$, $p = 0.076 > 0.05$); therefore, the perception of informational fairness did not differ between the different transparency levels, and the provision of informational transparency did not lead to a higher perception of informational fairness of AM. However, since the *p*-value was only slightly above the significance level – and we were therefore able to recognize a trend toward significance – we conducted a Mann–Whitney U test to test potential differences among the four groups. The results showed that there was no significant difference between Group 1 (i.e., control group with no transparency) ($M = 3.17$) and Group 3 (i.e., informational transparency) ($M = 3.19$), suggesting that there was no difference regarding the perception of informational fairness between the participants who received no transparency and those who received informational transparency ($F = 1741.5$, $p = 0.996 > 0.05$). Since we did not identify statistically significant differences between providing no transparency and providing informational transparency in the perception of informational fairness, we did not find support for H1b. However, our results indicated that there were significant differences between Group 1 (i.e., control group with no transparency) ($M = 3.17$) and Group 4 (i.e., both distributive and informational transparency) ($M = 3.61$, $F = 1384.5$, $p = 0.029 < 0.05$). This indicates that the participants who

received both distributive and informational transparency perceived informational fairness as significantly different from those who did not receive any transparency. Once again, to further determine the mean differences between the groups, we calculated the effect sizes according to Cohen's d (Cohen 1988), and the effect size between Groups 1 and 4 (i.e., no transparency and both distributive and informational transparency) ($d = -0.43$) indicated a small to moderate effect. These results highlight that, on average, the control group (i.e., no transparency) had a lower perception of informational fairness than the group that received both distributive and informational transparency.

To analyze whether providing **algorithmic transparency** on DLPs as a combination leads to a higher perception of fairness of AM (H1c), we calculated an additional Kruskal–Wallis test. Our results showed that there were no significant differences in perceptions of distributive and informational fairness among the four independent groups (i.e., no transparency, distributive transparency, informational transparency, or both distributive and informational transparency) ($F = 0.621$, $p = 0.891 > 0.05$). Therefore, the perception of fairness did not differ between groups (i.e., transparency levels), and the provision of transparency in general did not lead to a higher overall perception of fairness of AM. Since we did not identify statistically significant differences between providing no transparency and providing transparency, we did not find support for H1c.

4.1.2 The Impact of Trust in AM and Affinity for Technology on the Relationship between Transparency and Fairness

To determine the extent to which affinity for and trust in AM impact the relationship between fairness and transparency, we conducted a moderation analysis. Because we had both normally distributed and nonnormally distributed samples, we chose a nonparametric test for this analysis. Moreover, since residuals were not normally distributed, we conducted a moderation analysis using the generalized linear model (GAMLj) in Jamovi statistical software.⁸ This method allowed us to estimate the interaction effects between transparency and the moderator variables to test whether the relationship between transparency and fairness was influenced by those moderators.

First, we analyzed whether affinity for technology moderated the effect of transparency on the perception of fairness of AM. Our results revealed that affinity for technology did not have a moderating effect on the relationship between algorithmic transparency and fairness

perception since no significant effect could be found ($p = 0.842 > 0.05$). Due to the lack of a significant interaction effect, we rejected H2a. To further investigate whether affinity for technology had a moderating effect on the relationship between distributive transparency and fairness or between informational transparency and fairness, we conducted additional moderation analyses. The results indicated that affinity for technology moderated neither the effect on the relationship between distributive transparency and distributive fairness perception ($p = 0.359 > 0.05$) nor the effect on the relationship between informational transparency and of informational fairness perception ($p = 0.284 > 0.05$), since no significant effects could be identified. Due to the lack of significant interaction effects, we rejected both H2b and H2c.

Second, we analyzed whether trust in AM had a moderating effect on the relationship between transparency and the perception of fairness of AM. Our results showed that trust in AM did not have a moderating effect on the relationship between algorithmic transparency and fairness perception, since no significant effect could be identified ($p = 0.520 > 0.05$). Due to the lack of a significant interaction effect, we therefore also rejected H3a. In the next step, we further analyzed whether trust in AM had a moderating effect on the relationship between distributive transparency and fairness and between informational transparency and fairness. Our results revealed that trust in AM moderated neither the effect on the relationship between distributive transparency and the distributive fairness perception ($p = 0.190 > 0.05$) nor the effect on the relationship between informational transparency and informational fairness perception ($p = 0.818 > 0.05$), since no significant effects could be identified. Therefore, we also had to reject H3b and H3c.

In addition, we constructed an interaction term from affinity for and trust in AM and analyzed whether this interaction effect had a moderating role on any fairness dimension. However, we found that the interaction between affinity and trust was not a significant moderator for perceived ($p = 0.191 > 0.05$), distributive ($p = 0.156 > 0.05$), or informational ($p = 0.578 > 0.05$) fairness. Thus, our results suggest that affinity for technology, trust in AM, and their interaction had no moderating effect on the interaction between the different types of transparency and the associated perceptions of fairness of AM on a DLP.

4.1.3 Additional Findings

We conducted additional exploratory analyses to gain deeper insights into our findings beyond our literature-informed hypotheses. Through this exploratory research, we intended to shed light on the relationships among the

⁸ <https://www.jamovi.org>.

constructs of our hypotheses, which are not backed by the literature but may generate additional explanations. We considered this useful because our literature-driven hypothesis analysis did not yield many significant explanations, prompting us to shed further light on fairness. First, Spearman's correlation tests were conducted to gain more insight into the relations between humans' trust in computers (from the HCTS) – that is, trust in AM – or their affinity for technology interaction (from the ATI scale) and perceived fairness, informational fairness, and distributive fairness. A high correlation between participants' trust in AM and perceived fairness was identified ($\rho = 0.566$, $p < 0.001$). Participants who had generally high trust in AM also had high perceived fairness of AM on the DLP. A second correlation was identified between the participants' affinity for technology interaction and perceived fairness ($\rho = 0.311$, $p < 0.001$). Those who reported a high affinity for technology interaction also reported high perceived fairness. Another moderate correlation was identified between participants' trust in AM ($\rho = 0.413$, $p < 0.001$) and affinity for technology ($\rho = 0.373$, $p < 0.001$) and their perceived informational fairness. Participants who reported a high level of trust in AM and an affinity for technology in general also showed a high level of informational fairness. Lastly, our explorative results revealed a low correlation between participants' trust in AM ($\rho = 0.268$, $p < 0.001$) and affinity for technology ($\rho = 0.225$, $p < 0.001$) and their perceived distributive fairness. Therefore, participants who revealed a high level of trust in AM and an affinity for technology in general also showed a high level of distributive fairness.

Another interesting finding was the tendency for women ($M = 4.94$), compared to men ($M = 4.79$), to perceive a moderately stronger sense of fairness related to AM activities. This was also shown for distributive fairness (women: $M = 3.73$, men: $M = 3.59$) and informational fairness (women: $M = 3.41$, men: $M = 3.25$), with women showing moderately higher fairness than men. Moreover, participants employed full-time had the highest perception of AM fairness ($M = 5.89$), while those who were seeking work, in general, had the lowest levels of perceived fairness ($M = 4.50$), distributive fairness ($M = 3.52$), and informational fairness ($M = 3.08$). Considering the influence of the participants' frequency of working as digital workers, it was observed that those who worked daily tended to disclose the highest levels of perceived fairness ($M = 5.01$) and distributive fairness ($M = 3.73$), while those who worked on an annual basis perceived informational fairness ($M = 3.69$) as the highest. Participants who worked annually as digital workers showed the lowest level of perceived fairness ($M = 4.61$), and those who worked monthly as digital workers revealed the lowest level of

distributive fairness ($M = 3.49$) and informational fairness ($M = 2.89$).

4.2 Qualitative Findings

To gather additional insights into the participants' perceptions, we asked three open-ended questions. Depending on whether the participants found the allocation of studies fair or unfair, we asked, “*Why did you perceive the study allocation to be fair/unfair?*” Similarly, due to our interest in whether they rejected any study, we asked, “*Why did/didn't you reject the study/studies?*” Lastly, we asked for the participants' perceptions of the points calculation, asking, “*Why did/didn't you perceive the points for the studies calculated by the algorithm to be fair?*” All participants gave short answers to these questions. To gain an understanding of the participants' perceptions, we conducted an inductive qualitative content analysis according to Mayring (2015), with the relevant findings reported in the following sections.

4.2.1 Fairness Perception

Across all conditions, the most common reason participants found the study allocation fair was the perception that personal preferences were considered in the process. For example, “It allocates studies equally based on the information you provided to the algorithms” (Participant 11, G1) and “It matched my profile and still was able to offer variety to choose and did not force the task to be completed” (Participant 137, G3). Furthermore, another reason why the studies' assignment was suitable for the groups provided with information (Groups 2, 3, and 4) was that this was done objectively. For example, “Algorithms are designed with fairness in mind and are based on high-quality, representative data” (Participant 97, G2) and “There wasn't any inherent bias in the algorithm; it assigned it according to need” (Participant 202, G4). These answers indicate that the participants who were provided with transparency did not question the objectivity of algorithmic decisions on the DLP. This distinguishes them from the group that did not receive transparency, whose participants noticed a lack of information and often claimed that they simply had no reason to think the allocation was unfair. For example, “I don't have a lot of information to decide if there was any bias regarding study allocation – assuming it was assigned based on fair criteria” (Participant 13, G1) and “Because it didn't provide me with how and why tasks were assigned” (Participant 29, G1).

In summary, although the participants in all groups perceived the study allocation as fair, the responses from

Group 1 indicated a lack of transparency (i.e., the control group with no transparency).

4.2.2 Rejections

A consistent but small number of participants rejected studies in each group – 6, 19, 6, and 6 participants in Groups 1–4, respectively. Most participants found the studies fair – for example, “Because [they] trust the algorithm” (e.g., Participant 15, G1). Others rejected studies because they did not have enough time – for example, “I rejected some studies simply because of time on my part” (Participant 92, G2) – or they were not interested in the subject, such as “Because it does not align with my interest” (Participant 50, G1)).

Moreover, the participants who did not reject any studies did so because they were curious about them and felt pleasure, stating that they were “interested in exploring all of the studies” (Participant 19, G1) and “enjoy them” (Participant 99, G2). This view was held by the participants in all groups. Additional reasons were that they tried to get the full experience and, therefore, wanted to test the algorithm – for example, “because of the experience and to test AM’s capabilities and fairness” (Participant 97, G2). In line with perceived fairness, all the groups’ participants found that completing the studies did not cause negative consequences.

In summary, most of the participants did not reject any studies, since they liked to complete them and found them suitable. Those who rejected studies justified doing so because they did not align with their interests or had concerns regarding time.

4.2.3 Fair Calculations

Similar to perceived AM fairness, participants were indifferent about the fairness of the reward calculations across all groups. However, we observed two positions in each group. On the one hand, the participants reasoned that there was a mismatch between the time required and the effort they had to put in to complete a study. For example, “While some are fair, others require too much time and effort” (Participant 8, G1). Additionally, some of the participants in Groups 1 and 3 did not understand how points were allocated or how they were rewarded. For example, “It was not clear enough how much one would be rewarded” (Participant 121, G3). Those with an opposite view argued that the number of points was appropriate for the studies and that they perceived them to be fair, “since the criteria for scoring points were clear and understandable before the start” (Participant 68, G2). More interestingly, we found claims about a lack of transparency more frequently in Group 4 – for example, “I don’t know how the

algorithm works to calculate scores, but I think it was fair” (Participant 188, G4) – indicating that to be perceived as fair, it is insufficient to prove only the number of points compared to the time duration.

In summary, the participants’ answers indicated that more information did not necessarily increase the perceived fairness of the calculations, whereas mere truthfulness about the rewards led to perceived fairness; no information about the calculations was provided.

5 Discussion

We found that the provision of transparency in the underlying processes of a DLP did not influence workers’ perception of fairness compared to a DLP with no transparency regarding its functions and algorithms. Only distributive transparency had a small but significant effect on distributive fairness. Furthermore, we observed that neither affinity for technology nor trust in AM moderated the effect of transparency on the perception of fairness. However, affinity and trust were immediately correlated with all dimensions of fairness. Qualitative responses explained that the participants did not question fairness but that trust in the algorithm led them to perceive it as fair. Further, participants described the algorithm as unbiased and objective, with their main reason for accepting or rejecting a task being individual interest or lack thereof.

In contrast to extant work (e.g., Benlian et al. 2022; Cameron et al. 2023), our results do not suggest that providing transparency regarding AM decision-making processes and practices increases workers’ perceptions of fairness, even when providing them with information about the algorithm’s process of assigning studies and rewards. Instead, they challenge previous findings regarding perceptions of fairness and transparency (Parent-Rocheleau and Parker 2022; Zhang et al. 2022). The participants explained that they did not question the algorithm’s fairness but considered it unbiased and objective. Even those who had no information about how the AM worked – and thus experienced the least amount of transparency compared to the treatment groups – perceived the algorithm’s task allocation as fair, thereby corroborating workers’ general trust in algorithms’ rationality (Bao et al. 2021).

Algorithms are often considered inscrutable due to their inherent complexity – so-called black boxes that hinder users from interpreting outcomes (Kim et al. 2020; Maedche et al. 2019; Wanner 2021; Wanner et al. 2022). The inability to scrutinize algorithmic decisions makes it impossible for individuals to assess fairness (Heßler et al. 2022). Workers on a DLP do not question or condemn the algorithm’s decisions until they understand how it acts (Ochmann et al. 2021; Wanner 2021). Our findings extend

this understanding by demonstrating that distributive and informational transparency do not mitigate this issue.

Our results suggest that informational transparency – that is, providing participants with information about the process of assigning studies – did not have an impact on the perceived informational fairness of AM on a DLP, contrary to the findings of earlier research (e.g., Maedche et al. 2019; Parent-Rochelleau and Parker 2022; Peters et al. 2020; Pezzo et al. 2022). While research has shown that transparency in a system’s design fosters human–machine interaction (Berger et al. 2021), we did not find a similar effect regarding the perception of informational fairness in AM. This could be explained by the fact that participants who received information related to the functionalities and decision-making of the algorithm engaged more intensively with the allocation process and consequently had a weakened perception of fairness. First, this aligns with workers’ aversion to being guided and managed by an algorithm (Dietvorst et al. 2015; 2018), which may be weighted higher than information about the fairness of an algorithm and the accuracy of algorithmic recommendations. Second, individuals are more likely to question algorithmic recommendations when provided with information in this way (Ochmann et al. 2021; Wanner 2021). With respect to our research question, our results suggest that users who do not know how and why an algorithm has made a decision are more likely to accept its decisions and thus do not perceive particular unfairness, as they do not question the algorithm’s actions and decisions. The qualitative responses corroborated this, with participants stating they had no grounds to question the algorithm in the nontransparency condition.

As hypothesized, our results showed that distributive transparency – that is, providing participants with information about the rewards compared to their respective contributions – increases the perceived distributive fairness of AM on a DLP. Hence, we corroborate previous findings that consider algorithmic transparency to positively influence perceived distributive fairness (e.g., Bujold et al. 2022; Jabagi et al. 2024). Since we observed that distributive transparency had an effect, while informational transparency did not, we assume that the assessment of fairness depends on how the AM operates in relation to workers’ evaluations and rewards rather than on the processes that lead to the allocation of tasks. Put differently, the DLP workers appeared to evaluate fairness in terms of payment instead of awareness of algorithmic processes. This aligns with growing concerns regarding worker dehumanization and the exploitation of working conditions (Zhang et al. 2025). In this respect, workers are more concerned about getting paid than about the fair allocation of jobs, which highlights their precarious conditions (Chan 2022; Zhang et al. 2025). Thus, our results transfer these

already-known conditions from the gig economy literature (e.g., Hu and Han 2021) to work on DLPs.

Contrary to our initial hypotheses, our findings do not suggest that trust in AM or affinity for technology moderates the effect of transparency on fairness perception. Therefore, it appears that fairness was perceived solely as a result of the perceived transparency type. However, since previous research has found a link between trust and perceptions of fairness, we also tested whether trust correlates with perceptions of AM fairness regardless of transparency type (Höddinghaus et al. 2021; Krasnova et al. 2014; Spiekermann et al. 2022). We contribute to this understanding through our finding that a high level of perceived fairness is likely associated with high trust in AM and affinity for technology. This indicates that – regardless of the provision of transparency on DLP – workers’ affinity for technology and trust in AM have a positive impact on perceptions of AM fairness. Therefore, we propose that these personal attributes may be descriptive – rather than revealing – for evaluating AM fairness, such as information about AM’s functionalities and procedures (i.e., informational transparency). Moreover, the participants explained that they perceived the algorithm as fair because they had trust in it. This suggests that users perceive an algorithm on DLP as fairer if they generally have a high level of trust in the algorithm and understand how it works.

5.1 Contribution to Research

Our findings contribute to information systems research by providing knowledge on the perception of AM on DLPs, which constitute a fast-growing, technology-driven business model for organizations. Our results suggest that transparency in AM does not necessarily lead to increased perceived fairness, although it has been suggested as a means to overcome fairness issues (Benlian et al. 2022; Cameron et al. 2023; Cheng and Foley 2019; Robert et al. 2020). We showed that even if informational transparency is provided, more information is demanded, which can hardly be provided due to the black box characteristics of algorithmic systems. The high complexity of algorithms is difficult to explain to human workers. Thus, more transparency will likely risk mental exhaustion (Yang et al. 2024).

Consequently, this research aligns with the research stream that advocates the limited ability of transparency to increase perceived fairness (Ochmann et al. 2021; Wanner 2021). Specifically, our findings provide knowledge on the role of informational transparency when workers interact with AM on DLPs. Our findings revealed no empirical differences in the perception of fairness across different types of transparency. On the contrary, the provision of information could challenge workers to question an

algorithm. Hence, we oppose research that suggests that informational transparency facilitates fairness (e.g., Parent-Rocheleau and Parker 2022; Zhang et al. 2022).

Furthermore, we contribute to the OJT concept of perceived fairness in the context of AM. While organizational justice research asserts that fairness perceptions are grounded in rational judgments (Ganegoda et al. 2015), our findings explain that the perception of being treated fairly is subjective to the worker. In fact, personal attitudes and tendencies toward technology – such as affinity (Franke et al. 2017) or trust (Gulati et al. 2019) – seem to be more important than precise information about functionalities and procedures. That is, the perception of the fairness of an algorithm is shaped by workers' stances toward the technology (e.g., thinking of algorithms as objective) rather than by the information provided. Thus, we propose that these factors are consistently accounted for in OJT. We invite future research to assess the end-to-end effects of how affinity and trust are forged for the perception of fairness.

In summary, this study enhances the overall well-being of digital workers in digital businesses by investigating a significant subject: fairness. In this respect, we highlight the insufficiency of algorithmic transparency to increase workers' perceived fairness of AM. Our work provides suggestions for other researchers conducting online studies. Specifically, it was the participants' interest in the studies' topics that was perceived as fair. In this respect, to maintain the fair treatment of participants, we recommend acquiring suitable and interested participants for online studies rather than focusing on the transparency of the experiment. We align with humanistic approaches to AM, emphasizing the need to overcome mechanistic AM (e.g., Cui et al. 2024). Further, we contribute to the algorithmic transparency literature – while Li et al. (2025) identified the negative effects of the transparent declaration of algorithms' inner processes, we also contribute to research by transferring the limited effect on OTJ's perceived fairness.

5.2 Contribution to Practice

Our findings also contribute to practice and are relevant to DLP developers and managers. Specifically, they suggest that an increase in informational transparency does not lead to workers' perception of AM as fairer. Consequently, DLP developers and managers must pay attention to not falling for *fairwashing* the working conditions of their platform. However, we do not suggest canceling all attempts to increase transparency on DLPs because the provision of distributive transparency has led to an increased perception of fairness. Rather, we argue that DLP workers should be surveyed to understand what they might find more meaningful with regard to which specific types of information

could increase their perceptions of fairness. By understanding workers' preferences, future approaches could be customized to DLP workers, their sector, and the situations they are in. This is especially important for workers who are averse to technology.

Our work also has implications for policymakers. In the finance sector, it is an established practice that algorithmic decision-making must be disclosed. The recent AI Act⁹ by the European Union poses wide-ranging legal requirements for AI applications in every context. It categorizes these applications in terms of their potential risks and requires various actions, of which transparency is one. In this respect, policymakers should consider the impact of transparency on the development of future legal acts. Specifically, they must update their instruments governing new technologies to enable fair working conditions in organizations with digital business models.

5.3 Limitations and Future Research

This study has some limitations. We analyzed the impact of transparency on the perceived fairness of workers on a DLP. Therefore, the results should be interpreted carefully, as they might not be generalizable or transferable to other contexts. Since we conducted our study only with digital workers on the Prolific platform, this could be another limitation of the study. It is conceivable that creating transparency by providing information on AM functionalities and decision-making processes would have a different effect on other platforms whose focus is not based on study allocation, such as DLPs including Uber. Another limitation is the use of scales with mixed Likert anchors (i.e., 5, 6, and 7-point Likert anchors), which may have affected comparability across measures and their interpretation as a metric scale. Moreover, this study examined the impact of transparency on perceptions of fairness only among English-speaking workers. This limitation restricted our sample to a specific group of people, but considering additional cultural backgrounds would have extended the applicability of the results. While our definition of algorithmic transparency encompassed monitoring and controlling DLP workers, it did not consider the ability of systems to interact with workers or integrate their feedback, thereby limiting the scope to noninteractive forms of transparency.

To address these limitations, future research should examine how algorithmic systems can enable not merely transparency but effective interaction, allowing workers to question, shape, or influence algorithmic processes through responsive feedback mechanisms. Moreover, it should explore the effect of transparency through the provision of

⁹ <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX:52021PC0206>.

information on AM functionalities and decisions on users' perceptions of fairness in other DLP contexts to investigate whether differences in outcomes can be identified. We also propose evaluating contextual differences. Specifically, we recommend examining whether providing information in other settings – such as ride-hailing versus digital work – leads to perceptions of fairness or unfairness among users, thus retesting the effect of transparency on fairness. Given our findings that fairness perceptions are subject to individual evaluations, the emergence of cognitive processes that constitute the evaluation of perceived fairness is useful for examining the human drive for sense-making (cf. Chater and Loewenstein 2016). Since our results are specific to DLPs – and the use of technology is becoming more common in other organizational forms (Jarrahi et al. 2021) – future research should also investigate the impact of transparency on fairness perceptions of AM in traditional firms. For the IS community, this means that further research is needed to determine why transparency within AM does not help create fairness in human–machine interactions. In particular, research should investigate how the use of AM can lead to greater fairness among digital workers.

Finally, since we included scientific terms, such as algorithmic matching and controlling, in our study design, this could have potentially had a negative impact on the participants' understanding and, thus, our results. To address this limitation, future research is invited to make the study design more accessible to nonacademic audiences.

6 Conclusion

We studied the influence of algorithmic transparency on the perception of fairness in AM on a DLP. To this end, we conducted an online experiment featuring three transparency conditions – distributive transparency, informational transparency, and a combination of distributive and informational transparency – and a control group. Correspondingly, we measured perceived, informational, and distributive fairness as dependent variables in response to the experimental conditions, moderated by affinity for technology and trust in AM. Our findings suggest that transparency hardly influences any fairness dimensions. Further exploration revealed that individual evaluations, such as affinity and trust in AM, significantly correlated with participants' perceptions of fairness, suggesting that their personal stance may be a better explanation for fairness perception. An additional qualitative inquiry corroborated this notion, as the participants explained that they appreciated the fairness of AM due to its objectivity.

Hence, we suggest that IS scholarship moves forward by including individual evaluations to approach fairer AM.

Supplementary Information The online version contains supplementary material available at <https://doi.org/10.1007/s12599-025-00963-1>.

Funding Open Access funding enabled and organized by Projekt DEAL.

Open Access This article is licensed under a Creative Commons Attribution 4.0 International License, which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons licence, and indicate if changes were made. The images or other third party material in this article are included in the article's Creative Commons licence, unless indicated otherwise in a credit line to the material. If material is not included in the article's Creative Commons licence and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder. To view a copy of this licence, visit <http://creativecommons.org/licenses/by/4.0/>.

References

- Alexander S, Ruderman M (1987) The role of procedural and distributive justice in organizational behavior. *Soc Justice Res* 1(2):177–198. <https://doi.org/10.1007/BF01048015>
- Bao Y, Cheng X, De Vreede T, De Vreede GJ (2021) Investigating the relationship between AI and trust in human-AI collaboration. In: Proceedings of the annual Hawaii international conference on system sciences, Kauai, pp 607–616. <https://doi.org/10.24251/HICSS.2021.074>
- Barredo Arrieta A, Díaz-Rodríguez N, Del Ser J, Bennetot A, Tabik S, Barbado A, García S, Gil-Lopez S, Molina D, Benjamins R, Chatila R, Herrera F (2020) Explainable artificial intelligence (XAI): concepts, taxonomies, opportunities and challenges toward responsible AI. *Inf Fus* 58:82–115. <https://doi.org/10.1016/j.inffus.2019.12.012>
- Benlian A, Wiener M, Cram WA, Krasnova H, Maedche A, Möhlmann M, Recker J, Remus U (2022) Algorithmic management: bright and dark sides, practical implications, and research opportunities. *Bus Inf Syst Eng* 64(6):825–839. <https://doi.org/10.1007/s12599-022-00764-w>
- Berger B, Adam M, Rühr A, Benlian A (2021) Watch me improve: algorithm aversion and demonstrating the ability to learn. *Bus Inf Syst Eng* 63(1):55–68. <https://doi.org/10.1007/s12599-020-00678-5>
- Beugre CD, Baron RA (2001) Perceptions of systemic justice: the effects of distributive, procedural, and interactional justice. *J Appl Soc Psychol* 31(2):324–339. <https://doi.org/10.1111/j.1559-1816.2001.tb00199.x>
- Bitzer T, Wiener M, Cram WA (2023) Algorithmic transparency: concepts, antecedents, and consequences: a review and research framework. *Commun Assoc Inf Syst* 52(1):293–331. <https://doi.org/10.17705/1CAIS.05214>
- Bujold A, Parent-Rochelleau X, Gaudet MC (2022) Opacity behind the wheel: the relationship between transparency of algorithmic management, justice perception, and intention to quit among

- truck drivers. *Comput Hum Behav Rep* 8:100245. <https://doi.org/10.1016/j.chbr.2022.100245>
- Burrell J (2016) How the machine ‘thinks’: understanding opacity in machine learning algorithms. *Big Data Soc* 3(1):205395171562251. <https://doi.org/10.1177/2053951715622512>
- Cameron L, Lamers L, Leicht-Deobald U, Lutz C, Meijerink J, Möhlmann M (2023) Algorithmic management: Its implications for information systems research. *Commun Assoc Inf Syst* 52(1):556–574. <https://doi.org/10.17705/1CAIS.05221>
- Chan NK (2022) Algorithmic precarity and metric power: Managing the affective measures and customers in the gig economy. *Big Data Soc* 9(2):20539517221133780. <https://doi.org/10.1177/20539517221133779>
- Chan J, Wang J (2018) Hiring preferences in online labor markets: evidence of a female hiring bias. *Manag Sci* 64(7):2973–2994. <https://doi.org/10.1287/mnsc.2017.2756>
- Chater N, Loewenstein G (2016) The under-appreciated drive for sense-making. *J Econ Behav Organ* 126:137–154. <https://doi.org/10.1016/j.jebo.2015.10.016>
- Cheng M, Foley C (2019) Algorithmic management: the case of Airbnb. *Int J Hosp Manag* 83:33–36. <https://doi.org/10.1016/j.ijhm.2019.04.009>
- Cohen J (1988) *Statistical power analysis for the behavioral sciences*, 2nd edn. Routledge Lawrence Erlbaum, Hillsdale. <https://doi.org/10.4324/9780203771587>
- Cohen-Charash Y, Spector PE (2001) The role of justice in organizations: a meta-analysis. *Organ Behav Hum Decis Process* 86(2):278–321. <https://doi.org/10.1006/obhd.2001.2958>
- Colquitt JA (2001) On the dimensionality of organizational justice: a construct validation of a measure. *J Appl Psychol* 86(3):386–400. <https://doi.org/10.1037/0021-9010.86.3.386>
- Colquitt JA, Wesson MJ, Porter COLH, Conlon DE, Ng KY (2001) Justice at the millennium: a meta-analytic review of 25 years of organizational justice research. *J Appl Psychol* 86(3):425–445. <https://doi.org/10.1037/0021-9010.86.3.425>
- Colquitt JA, Scott BA, Rodell JB, Long DM, Zapata CP, Conlon DE, Wesson MJ (2013) Justice at the millennium, a decade later: a meta-analytic test of social exchange and affect-based perspectives. *J Appl Psychol* 98(2):199–236. <https://doi.org/10.1037/a0031757>
- Cram WA, Wiener M, Tarafdar M, Benlian A (2022) Examining the impact of algorithmic control on Uber drivers’ technostress. *J Manag Inf Syst* 39(2):426–453. <https://doi.org/10.1080/07421222.2022.2063556>
- Cui T, Li S, Chen K, Bailey J, Liu F (2023) Designing fair AI systems: exploring the interaction of explainable AI and task objectivity on users’ fairness perception. In: Pacific Asia conference on information systems proceedings, Nanchang. <https://aisel.aisnet.org/pacis2023/161>
- Cui T, Tan B, Shi Y (2024) Fostering humanistic algorithmic management: A process of enacting human-algorithm complementarity. *J Strateg Inf Syst* 33(2):101838. <https://doi.org/10.1016/j.jsis.2024.101838>
- Diakopoulos N, Koliska M (2017) Algorithmic transparency in the news media. *Digit Journal* 5(7):809–828. <https://doi.org/10.1080/21670811.2016.1208053>
- Dietvorst BJ, Simmons JP, Massey C (2015) Algorithm aversion: people erroneously avoid algorithms after seeing them err. *J Exp Psychol Gen* 144(1):114–126. <https://doi.org/10.1037/xge0000033>
- Dietvorst BJ, Simmons JP, Massey C (2018) Overcoming algorithm aversion: people will use imperfect algorithms if they can (even slightly) modify them. *Manag Sci* 64(3):1155–1170. <https://doi.org/10.1287/mnsc.2016.2643>
- Douglas BD, Ewell PJ, Brauer M (2023) Data quality in online human-subjects research: comparisons between MTurk, prolific, cloudresearch, qualtrics, and SONA. *PLoS One* 18(3):e0279720. <https://doi.org/10.1371/journal.pone.0279720>
- Dunn M, Munoz I, Jarrahi MH (2023) Dynamics of flexible work and digital platforms: task and spatial flexibility in the platform economy. *Digit Bus* 3(1):100052. <https://doi.org/10.1016/j.digbus.2022.100052>
- Faraj S, Pachidi S, Sayegh K (2018) Working and organizing in the age of the learning algorithm. *Inf Organ* 28(1):62–70. <https://doi.org/10.1016/J.INFOANDORG.2018.02.005>
- Felzmann H, Villaronga EF, Lutz C, Tamò-Larriex A (2019) Transparency you can trust: Transparency requirements for artificial intelligence between legal norms and contextual concerns. *Big Data Soc*. <https://doi.org/10.1177/2053951719860542>
- Franke T, Attig C, Wessel D (2017) Assessing affinity for technology interaction: the affinity for technology interaction (ATI) scale. <https://doi.org/10.13140/RG.2.2.28679.50081>
- Gal U, Jensen TB, Stein MK (2020) Breaking the vicious cycle of algorithmic management: a virtue ethics approach to people analytics. *Inf Organ* 30(2):100301. <https://doi.org/10.1016/j.infoandorg.2020.100301>
- Gal U, Jensen TB, Stein MK (2017) People analytics in the age of big data: an agenda for IS research. In: International conference on information systems proceedings, Seoul. <http://aisel.aisnet.org/icis2017/TransformingSociety/Presentations/1>
- Ganegoda DB, Folger R (2015) Framing effects in justice perceptions: prospect theory and counterfactuals. *Organ Behav Hum Decis Process* 126:27–36. <https://doi.org/10.1016/j.obhdp.2014.10.002>
- Geissinger A, Laurell C, Öberg C, Sandström C, Suseno Y (2022) The sharing economy and the transformation of work: evidence from Foodora. *Pers Rev* 51(2):584–602. <https://doi.org/10.1108/PR-08-2019-0450>
- Greenberg J (1990) Organizational justice: yesterday, today, and tomorrow. *J Manag* 16(2):399–432. <https://doi.org/10.1177/014920639001600208>
- Gulati SN, Sousa SC, Lamas D (2019) Design, development and evaluation of a human-computer trust scale. *Behav Inf Technol* 38(10):1004–1015. <https://doi.org/10.1080/0144929X.2019.1656779>
- Heilala V, Kelly R, Saarela M, Jääskelä P, Kärkkäinen T (2023) The finnish version of the affinity for technology interaction (ATI) scale: psychometric properties and an examination of gender differences. *Int J Hum-Comput Interact* 39(4):874–892. <https://doi.org/10.1080/10447318.2022.2049142>
- Heinrich K, Vu MA, Vysochyna A (2022) Algorithms as a manager: a critical literature review of algorithm management. In: International conference on information systems proceedings, Copenhagen. https://aisel.aisnet.org/icis2022/is_futureofwork/is_futureofwork/9
- Heßler PO, Pfeiffer J, Hafenbrädl S (2022) When self-humanization leads to algorithm aversion: what users want from decision support systems on prosocial microlending platforms. *Bus Inf Syst Eng* 64(3):275–292. <https://doi.org/10.1007/s12599-022-00754-y>
- Höddinghaus M, Sondern D, Hertel G (2021) The automation of leadership functions: Would people trust decision algorithms? *Comput Hum Behav* 116:106635. <https://doi.org/10.1016/j.chb.2020.106635>
- Hofeditz L, Clausen S, Rieß A, Mirbabaie M, Stieglitz S (2022) Applying XAI to an AI-based system for candidate management to mitigate bias and discrimination in hiring. *Electron Mark* 32:2207–2233. <https://doi.org/10.1007/s12525-022-00600-9>

- Hu B, Han S (2021) Distributive justice: investigating the impact of resource focus and resource valence. *J Bus Psychol* 36(2):225–252. <https://doi.org/10.1007/s10869-019-09668-1>
- Jabagi N, Croteau AM, Audebrand LK, Marsan J (2019) Gig-workers' motivation: thinking beyond carrots and sticks. *J Manag Psychol* 34(4):192–213. <https://doi.org/10.1108/JMP-06-2018-0255>
- Jabagi N, Croteau AM, Audebrand LK, Marsan J (2024) Fairness in algorithmic management: bringing platform-workers into the fold. In: Hawaii international conference on system sciences proceedings. https://aisel.aisnet.org/hicss-57/cl/ai_and_future_work/5
- Jarrah MH, Newlands G, Lee MK, Wolf CT, Kinder E, Sutherland W (2021) Algorithmic management in a work context. *Big Data Soc.* <https://doi.org/10.1177/20539517211020332>
- Jiang LD, Ravichandran T, Kuruzovich J (2023) Review moderation transparency and online reviews: evidence from a natural experiment. *MIS Q* 47(4):1693–1708. <https://doi.org/10.25300/MISQ/2023/16216>
- Kässi O, Lehdonvirta V, Stephany F (2021) How many online workers are there in the world? A data-driven assessment [version 3; peer review: 4 approved]. *Open Res Eur* 1:53. <https://doi.org/10.12688/openreseurope.13639.3>
- Kellogg KC, Valentine MA, Christin A (2020) Algorithms at work: the new contested terrain of control. *Acad Manag Ann* 14(1):366–410. <https://doi.org/10.5465/annals.2018.0174>
- Kim B, Park J, Suh J (2020) Transparency and accountability in AI decision support: explaining and visualizing convolutional neural networks for text information. *Decis Support Syst* 134:113302. <https://doi.org/10.1016/j.dss.2020.113302>
- Köchling A, Wehner MC (2020) Discriminated by an algorithm: a systematic review of discrimination and fairness by algorithmic decision-making in the context of HR recruitment and HR development. *Bus Res* 13(3):795–848. <https://doi.org/10.1007/s40685-020-00134-w>
- Kordzadeh N, Ghasemaghahi M (2022) Algorithmic bias: review, synthesis, and future research directions. *Eur J Inf Syst* 31(3):388–409. <https://doi.org/10.1080/0960085X.2021.1927212>
- Krämer J, Wiewiorra L (2015) When 'just' is just not enough: why consumers do not appreciate non-neutral internet access services. *Bus Inf Syst Eng* 57(5):325–338. <https://doi.org/10.1007/s12599-015-0398-9>
- Krasnova H, Veltri NF, El Garah W (2014) Effectiveness of justice-based measures in managing trust and privacy concerns on social networking sites: an intercultural perspective. *Commun Assoc Inf Syst* 35(1):83–108. <https://doi.org/10.17705/1cais.03504>
- Langer M, König CJ (2023) Introducing a multi-stakeholder perspective on opacity, transparency and strategies to reduce opacity in algorithm-based human resource management. *Hum Resour Manag Rev* 33(1):100881. <https://doi.org/10.1016/j.hrmr.2021.100881>
- Lee MK (2018) Understanding perception of algorithmic decisions: fairness, trust, and emotion in response to algorithmic management. *Big Data Soc* 5(1):1–16. <https://doi.org/10.1177/2053951718756684>
- Lee MK, Kusbit D, Metsky E, Dabbish L (2015) Working with machines: the impact of algorithmic and data-driven management on human workers. In: Conference on human factors in computing systems: proceedings, 1603–1612. <https://doi.org/10.1145/2702123.2702548>
- Li Y, Zhao L, Cao C, Yang D (2025) The double-edged sword effect of algorithmic transparency: an empirical study of gig workers' work disengagement under algorithmic management. *Inf Manag* 62(2):104100. <https://doi.org/10.1016/j.im.2025.104100>
- Lu X, Phang D, Ba S, Yao X (2024) The effects of featuring product sampling reviews on e-tailer websites. *J Assoc Inf Syst* 25(3):618–647. <https://doi.org/10.17705/1jais.00834>
- Lu Y, Gupta A, Ketter W, Heck EV (2014) Information transparency in multi-channel B2B auctions: a field experiment. In: International conference on information systems proceedings, Auckland. <https://aisel.aisnet.org/icsis2014/proceedings/EBusiness/42>
- Luo Y (2007) The independent and interactive roles of procedural, distributive, and interactional justice in strategic alliances. *Acad Manag J* 50(3):644–664. <https://doi.org/10.5465/AMJ.2007.25526452>
- Maedche A, Legner C, Benlian A, Berger B, Gimpel H, Hess T, Hinz O, Morana S, Söllner M (2019) AI-based digital assistants: opportunities, threats, and research perspectives. *Bus Inf Syst Eng* 61(4):535–544. <https://doi.org/10.1007/s12599-019-00600-8>
- Manyika J, Lund S, Bughin J, Robinson K, Mischke J, Mahajan D (2016) Independent work: choice, necessity, and the gig economy. McKinsey Global Institute. <https://www.mckinsey.com/~media/McKinsey/Featured%20Insights/Employment%20and%20Growth/Independent%20work%20Choice%20necessity%20and%20the%20gig%20economy/Independent-Work-Choice-necessity-and-the-gig-economy-Executive-Summary.pdf>. Accessed 20 Jul 2025
- Marjanovic O, Cecez-Kecmanovic D, Vidgen R (2022) Theorising algorithmic justice. *Eur J Inf Syst* 31(3):269–287. <https://doi.org/10.1080/0960085X.2021.1934130>
- Mayer RC, Davis JH, Schoorman FD (1995) An integrative model of organizational trust. *Acad Manag Rev* 20(3):709–734. <https://doi.org/10.2307/258792>
- Mayring P (2015) Qualitative content analysis: theoretical background and procedures. In: Bikner-Ahsbahs A et al (eds) Approaches to qualitative research in mathematics education. Springer, Dordrecht, pp 365–380. https://doi.org/10.1007/978-94-017-9181-6_13
- McKnight DH, Choudhury V, Kacmar C (2002) The impact of initial consumer trust on intentions to transact with a website: a trust building model. *J Strateg Inf Syst* 11(3–4):297–323. [https://doi.org/10.1016/S0963-8687\(02\)00020-3](https://doi.org/10.1016/S0963-8687(02)00020-3)
- McKnight DH, Carter M, Thatcher JB, Clay PF (2011) Trust in a specific technology: an investigation of its components and measures. *ACM Trans Manag Inf Syst* 2(2):1–25. <https://doi.org/10.1145/1985347.1985353>
- Mishra AK (1996) Organizational responses to crisis: The centrality of trust. In: Trust in organizations: frontiers of theory and research. Sage, New York, pp 261–287. <https://doi.org/10.4135/9781452243610.n13>
- Möhlmann M, Zalmanson L, Henfridsson O, Gregory RW (2021) Algorithmic management of work on online labor platforms: when matching meets control. *MIS Q* 45(4):1999–2022. <https://doi.org/10.25300/misq/2021/15333>
- Möhlmann M, Salge CAL, Marabelli M (2023) Algorithm sense-making: how platform workers make sense of algorithmic management. *J Assoc Inf Syst* 24(1):35–64. <https://doi.org/10.17705/1jais.00774>
- Möhlmann M (2021) Algorithmic nudges don't have to be unethical. *Harv Bus Rev* 22:1–6
- Morse L, Teodorescu MHM, Awwad Y, Kane GC (2022) Do the ends justify the means? Variation in the distributive and procedural fairness of machine learning algorithms. *J Bus Ethics* 181(4):1083–1095. <https://doi.org/10.1007/s10551-021-04939-5>
- Mosaferchi S, Califano R, Naddeo A (2023) How personality, demographics, and technology affinity affect trust in autonomous vehicles: a case study. *Hum Factors Transp* 95:227–236. <https://doi.org/10.54941/ahfe1003808>

- Moussawi S, Koufaris M, Benbunan-Fich R (2021) How perceptions of intelligence and anthropomorphism affect adoption of personal intelligent agents. *Electron Mark* 31(2):343–364. <https://doi.org/10.1007/s12525-020-00411-w>
- Ochmann J, Zilker S, Michels L, Tiefenbeck V, Laumer S (2021) The influence of algorithm aversion and anthropomorphic agent design on the acceptance of AI-based job recommendations. In: *Proceedings international conference on information systems*. https://aisel.aisnet.org/icis2020/is_workplace_fow/is_workplace_fow/4
- Parent-Rocheleau X, Parker SK (2022) Algorithms as work designers: how algorithmic management influences the design of jobs. *Hum Resour Manag Rev* 32(3):100838. <https://doi.org/10.1016/j.hrmr.2021.100838>
- Peters F, Pumplun L, Buxmann P (2020) Opening the black box: Consumer's willingness to pay for transparency of intelligent systems. In: *Proceedings of the 28th European conference on information systems, Marrakech, Paper 90*. https://aisel.aisnet.org/ecis2020_rp/90
- Pezzo MV, Nash BED, Vieux P, Foster-Grammer HW (2022) Effect of having, but not consulting, a computerized diagnostic aid. *Med Decis Making* 42(1):94–104. <https://doi.org/10.1177/0272989X211011160>
- Rani U, Furrer M (2021) Digital labour platforms and new forms of flexible work in developing countries: algorithmic management of work and workers. *Compet Change* 25(2):212–236. <https://doi.org/10.1177/1024529420905187>
- Robert LP, Pierce C, Marquis L, Kim S, Alahmad R (2020) Designing fair AI for managing employees in organizations: a review, critique, and design agenda. *Hum-Comput Interact* 35(5–6):545–575. <https://doi.org/10.1080/07370024.2020.1735391>
- Schmidt P, Biessmann F, Teubner T (2020) Transparency and trust in artificial intelligence systems. *J Decis Syst* 29(4):260–278. <https://doi.org/10.1080/12460125.2020.1819094>
- Schoonderwoerd TAJ, Van Zoelen EM, Van den Bosch K, Neerincx MA (2022) Design patterns for human-AI co-learning: a wizard-of-Oz evaluation in an urban-search-and-rescue task. *Int J Hum-Comput Stud* 164:102831. <https://doi.org/10.1016/j.ijhcs.2022.102831>
- Schulze L, Trenz M, Cai Z (2022) Algorithmic unfairness on digital labor platforms: how algorithmic management practices disadvantage workers. In: *Proceedings international conference on information systems, Copenhagen*. https://aisel.aisnet.org/icis2022/is_futureofwork/is_futureofwork/8
- Schulze L, Trenz M, Cai Z, Tan CW (2023) Fairness in algorithmic management: how practices promote fairness and redress unfairness on digital labor platforms. In: *Proceedings Hawaii international conference on system sciences*, 186–205. 10125/102652
- Sikayu SH, Rahmat M, Chan AN (2022) Fairness, transparency and attitude towards tax evasion amongst owners of SMEs. *Int J Serv Manag Sustain (IJSMS)* 7(1):185–206. <https://doi.org/10.24191/ijms.v7i1.17786>
- Sokol K, Flach P (2020) Explainability fact sheets. In: *Proceedings of the 2020 conference on fairness, accountability, and transparency*, 56–67. <https://doi.org/10.1145/3351095.3372870>
- Spiekermann S, Krasnova H, Hinz O, Baumann A, Benlian A, Gimpel H, Heimbach I, Köster A, Maedche A, Niehaves B, Risius M, Trenz M (2022) Values and ethics in information systems: a state-of-the-art analysis and avenues for future research. *Bus Inf Syst Eng* 64(2):247–264. <https://doi.org/10.1007/s12599-021-00734-8>
- Tarafdar M, Page X, Marabelli M (2022) Algorithms as co-workers: human algorithm role interactions in algorithmic work. *Inf Syst J* 33(2):232–267. <https://doi.org/10.1111/ISJ.12389>
- Wanner J, Herm LV, Heinrich K, Janiesch C (2022) The effect of transparency and trust on intelligent system acceptance: evidence from a user-based study. *Electron Mark* 32(4):2079–2102. <https://doi.org/10.1007/s12525-022-00593-5>
- Wanner J (2021) Do you really want to know why? Effects of AI-based DSS advice on human decisions. In: *Proceedings Americas conference on information systems, Montreal, 13:1–10*. https://aisel.aisnet.org/amcis2021/strategic_is/strategic_is/13
- Weiss A, Bernhaupt R, Schwaiger D, Altmaninger M, Buchner R, Tscheligi M (2009) User experience evaluation with a Wizard of Oz approach: technical and methodological considerations. In: *Proceedings 9th IEEE-RAS international conference on humanoid robots*, 303–308. <https://doi.org/10.1109/ICHR.2009.5379559>
- Wilson J, Rosenberg D (1988) Rapid prototyping for user interface design. In: Helander M (ed) *Handbook of human-computer interaction*. Elsevier, North-Holland, pp 859–875. <https://doi.org/10.1016/b978-0-444-70536-5.50044-0>
- Wood AJ, Graham M, Lehdonvirta V, Hjorth I (2019) Good gig, bad gig: autonomy and algorithmic control in the global gig economy. *Work Employ Soc* 33(1):56–75. <https://doi.org/10.1177/0950017018785616>
- Yang H, Li D, Hu P (2024) Decoding algorithm fatigue: the role of algorithmic literacy, information cocoons, and algorithmic opacity. *Technol Soc* 79:102749. <https://doi.org/10.1016/j.techsoc.2024.102749>
- Zhang S, Meng Z, Chen B, Yang X, Zhao X (2021) Motivation, social emotion, and the acceptance of artificial intelligence virtual assistants: trust-based mediating effects. *Front Psychol* 12:3441. <https://doi.org/10.3389/fpsyg.2021.728495>
- Zhang MM, Cooke FL, Ahlstrom D, McNeil N (2025) The rise of algorithmic management and implications for work and organisations. *New Technol Work Employ*. <https://doi.org/10.1111/ntwe.12343>
- Zhang A, Boltz A, Wang CW, Lee MK (2022) Algorithmic management reimagined for workers and by workers: centering worker well-being in gig work. In: *Proceedings conference on human factors in computing systems, New Orleans*. <https://doi.org/10.1145/3491102.3501866>
- Zhou YV, Leong C, Guo Z (2023) What is fair enough? Reconciling complementors' needs for fairness management on digital platforms. In: *proceedings international conference on information systems, Copenhagen*. https://aisel.aisnet.org/icis2023/sharing_econ/sharing_econ/11

Publisher's Note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.