



Research Article

Algorithm awareness: Why user awareness is critical for personal privacy in the adoption of algorithmic platforms?

Donghee Shin^{a,*}, Kerk F. Kee^b, Emily Y. Shin^{c,*}^a Zayed University, P.O. Box 144534, Abu Dhabi, United Arab Emirates^b Texas Tech University, 2500 Broadway, Lubbock, TX 79409, United States^c American Community School of Abu Dhabi, United Arab Emirates

ARTICLE INFO

Keywords:

Algorithm awareness
Self-efficacy
Personalized algorithms
Self-disclosure
Privacy calculus
Personal privacy

ABSTRACT

Understanding how algorithms shape users' online experiences is a prerequisite to developing an effective algorithm design. Due to the rapid algorithmification of platforms, it is timely to examine users' awareness of algorithms on online platforms because these algorithms can shape everyday decisions and interactions through mediating, gatekeeping, and structuring user interactions. Focusing on the role of algorithm awareness (AA) in the privacy calculus process, we investigate users' intention to disclose personal information when using a platform with personalized algorithms. By conceptualizing AA with a range of theoretical and behavioral variables, we examine how users' self-efficacy affects their privacy concerns when they adopt, consume, and interact with such platforms. The findings show that AA leads users to envisage, understand, and interact with algorithms depending on their understanding of the control of the information flow embedded within them. The awareness that users have regarding algorithms influences the trust of algorithmic processes and the way users evaluate privacy concerns and self-disclosures. The cognitive user processes of AA provide conceptual frameworks for algorithm design and a practical guideline for the design of personalized algorithms.

1. Introduction

As artificial intelligence (AI) becomes increasingly pervasive in society, algorithms have become more integrated into our lives, including when serving as gatekeepers of information gathering, content selection, and predictive analytics (Dwivedi et al., 2021). Algorithms have become the key mediator through which algorithmic curations shape individual, social, and economic life (Nishant, Kennedy, & Corbett, 2020). On social media, algorithms sort posts in users' stories or feeds (Schwartz & Mahnke, 2021). Through algorithms, news media applications prioritize what news users read in their feeds (Shin, Zhong, & Biocca, 2020). In entertainment, streaming platforms, such as *Spotify*, *Netflix*, and *Hulu*, use algorithms to drive their business, leading the way in providing personalized digital content for individual users (Siles, Segura-Castillo, Solís-Quesada, & Sancho, 2020).

Despite these notable benefits to users, the integration of algorithms for automated personalization raises ethical (Ashok, Madan, Joha, & Sivarajah, 2022) and privacy concerns (Lau, Zimmerman, & Schaub, 2018). Specifically, algorithms operate behind the interface, tracking user activities and regulating what becomes available to users without

users knowing what the algorithms are or their functions (Dwivedi et al., 2021). A notable concern is that algorithms do not simply facilitate the flow of content; they actively shape the content, and at times, they distort it with embedded biases (Akter et al., 2021; Gran, Booth, & Bucher, 2021). This concern raises a couple of key questions: "Can algorithms be explained to users as part of the open system?" "How can users be educated to better understand the operation of algorithmic personalization?" These questions justify the need to conceptualize algorithm awareness (AA) as an integral part of algorithm operations and practices (Eslami et al., 2015).

The increasing attention paid to AA is driven by the realization that users deserve to know how algorithms work (Hamilton, Karahalios, Sandvig, & Eslami, 2014), and user awareness should be a requirement in the algorithm codes of practice (Shin, 2021) to be consistent with General Data Protection Regulation (GDPR). The increasing importance of AA, inter alia, is the concrete means by which users become more informed about algorithmic systems (Kizilcec, 2016) because AA can influence users and their information sharing behaviors (Zarouali, Boerman, & Vreese, 2021). Thus, issues related to AA, such as fairness, explainability, accountability, and transparency (FEAT), have garnered

* Corresponding authors.

E-mail addresses: donghee.shin@zu.ac.ae (D. Shin), kerk.kee@ttu.edu (K.F. Kee), emilyshin@acs.sch.ae (E.Y. Shin).

tremendous public attention and have raised questions regarding how to operationalize and manage FEAT in AI development (Akteer et al., 2021; Diakopoulos, 2016). One such measure is to address FEAT by designing a user-controllable AI. As algorithms operate in a black-box like process, it is then important to examine in what ways users are aware of these issues, how they make sense of the issues, how they can meaningfully control their own interactions with AI by managing the data they choose to share and by evaluating their privacy and security practices, and what impacts FEAT has on user behaviors, particularly in their response to privacy concerns (Shin & Park, 2019). We argue that FEAT principles motivate AI professionals to take concrete steps toward designing responsible AI systems that facilitate meaningful human control and accountability.

In this article, we use a few key terms that deserve some clarifications upfront. First, although AI, algorithms, and platforms are intertwined in practice, they are conceptually distinct. We see AI as the overarching field that makes user personalization possible, however, algorithms are the specific mechanism that makes personalization possible through specific platforms (e.g., news platforms). Second, we use the term ‘algorithm’ and ‘algorithmic’ extensively throughout the article. When ‘algorithm’ is used with another noun, such as in the case of ‘algorithm design,’ we refer to the ‘design of, about, or related to the mechanism of algorithm’. When ‘algorithmic’ is used with a noun, such as in the case of ‘algorithmic design,’ we refer to a design that is driven and/or created by algorithms. In our title, we use the term ‘algorithm awareness’ to refer to user awareness of algorithm and how it works. Finally, once users understand how algorithms work and the potential harms created when the algorithmic platforms (and the companies behind them) have access to user’s personal information, the users may need to make an informed decision if and when to share such information. Because of the potential violation of information privacy that could in turn cause harms on the users, we decided to use ‘self-disclosure’ or ‘information disclosure’ to refer to users’ un/conscious sharing of their personal information with the platforms (and the companies behind them) when using such platforms. We realized that some readers may believe that ‘information sharing’ is the common term referring to the same behavior, but we believe ‘information disclosure’ is more descriptive of the behaviors in the context of risk.

Improving the awareness and understanding of algorithms and their capability among users have been a topic of debate. Shin (2021) clarifies how users experience algorithms and the extent to which they develop an understanding of algorithms. Eslami et al. (2015) argue that an AA design should provide a high degree of transparency to algorithmic processes to promote informed decisions regarding data sharing. Furthermore, Monzer, Moeller, Helberger, and Eskens (2020) propose that a user understanding of fairness, transparency, and accountability are necessary before users accept personalized algorithms. Understanding how users perceive an algorithm’s attributes, capability, recommendations, and quality of personalization is key to cultivating the awareness necessary to design and develop responsible AI (Zarouali et al., 2021).

Recent studies on AI adoption highlight the effects of FEAT on user subsequent behaviors, such as those related to privacy, trust, and intention to use (e.g., Chatterjee, Rana, Dwivedi, & Baabdullah, 2021; Fast & Jago, 2020; Gutierrez, O’Leary, Rana, Dwivedi, & Calle, 2019). Therefore, the causality between FEAT and privacy can be inferred in that the more relaxed users are about privacy—which then increases trust—the more users allow algorithms to collect their data. This proposition is consonant with the Privacy Calculus Theory (PCT; Culnan & Armstrong, 1999), which states that users’ intentions to disclose personal information are based on a risk-benefit analysis, weighing between the potential benefits and risks. Yet, it remains unknown how users perceive and understand FEAT, what constitutes user awareness, how it relates to privacy, and how trust facilitates data sharing. These questions are important in understanding personalized algorithms, given that they not only extend the knowledge of FEAT and privacy but also the relation

between the two concepts to further reveal the process of privacy calculus in sharing private information with personalized algorithms (Spanaki, Karafili, & Despoudi, 2021).

Against the backdrop of increasing concerns for privacy, it is important to identify the relationship between AA and FEAT by focusing on how users evaluate privacy in the context of algorithms through which users make sense of FEAT (Gran et al., 2021). The processes of establishing AA and evaluating privacy are necessary to facilitate users’ decisions to adopt algorithms (Siles et al., 2020). While considerable research has elaborated on individual factors, such as security and risk, there is a dearth of studies examining the role of privacy and user awareness as an explanatory factor to understand users’ decisions to share information with algorithms, indicating a clear research gap. We attempt to fill this gap by empirically examining how users perceive FEAT, how users develop their self-efficacy regarding algorithms, which roles AA plays in privacy decisions related to personalized algorithms, and how users decide when to self-disclose and share their personal information, especially in the context of AI-driven personalized/customized recommendations. The following research questions were developed to guide this study:

RQ1: What dimensions define the user awareness of algorithms, and how do these dimensions influence user experiences with algorithmically personalized services?

RQ2: How does AA influence users’ self-efficacy in the context of a personalized algorithmic process, and how does the perception influence users’ self-disclosure while using services with personalized algorithms?

RQ3: How does the user awareness of FEAT influence the privacy calculus toward the adoption of personalized algorithms?

The new implications of further understanding AA in relation to AI provide unique paths that are both practical and theoretical. Conceptualizing the cognitive process of AA and related constructs is the first attempt and advances contributions to the ongoing work on human–algorithm interactions (Shin, 2021). Although the effect of personalized algorithms on user behaviors has received much debate (Reisdorf & Blank, 2020), an extant empirical understanding of this effect, particularly from the user perspective of specific factors, such as privacy and self-disclosure, is still limited. To our knowledge, our work is the first to extend the existing theory on privacy valuation to user intention to disclose information through the theoretical development and empirical testing of the algorithm Privacy Calculus Model. The noble contribution of our research includes conceptual groundworks for AA, supporting the user privacy calculus process. Practically, the underlying role of AA in an algorithm lends strategic direction in developing human-centered AI against the dehumanizing trends of algorithmic operations (Borges, Laurindo, Spínola, Gonçalves, & Mattos, 2022; Swart, 2021). Our results provide operational criteria for an algorithm audit, making meaningful control over algorithms a challenge. User oversight based on AA is one of the requirements as a means of supporting and ensuring meaningful human control over algorithms by fulfilling core values advocated in ethical guidelines: fairness, transparency, and accountability. Our new results are also consonant with the increasing importance of user rights as proclaimed by the GDPR in the AI era. Our novel findings provide platform providers with guidance to better respond to GDPR and to improve their data privacy measures.

The remainder of this article is outlined as follows. Section 2 provides the conceptual background and related work on AA, privacy, and algorithms. Section 4 presents the data and method specifications. Afterward, Section 5 discusses the findings with the experiments, and Section 6 describes the theoretical and practical implications. Finally, Section 7 presents the final considerations.

2. Literature review

User awareness matters because it shapes behaviors (Gran et al., 2021) and provides a basis for the design of user-controllable AI. AA includes efficacy, which leads users to engage in appropriate privacy behaviors. In other words, AA generates particular ways of interacting with personalized algorithms. In this study, we examined how users understand the mechanisms of personalized algorithms, how they come to develop efficacy through awareness, and the implications of self-efficacy for evaluating privacy and self-disclosure. We broaden our understanding of how people make sense of algorithmic output in relation to data, results, and user knowledge by incorporating issues of privacy into the AA process.

2.1. Algorithm awareness

We conceptualize AA as similar to the case of ordinary people's knowledge about algorithmic systems, although the literature suggests its connections, including transparency, accountability, fairness (Shin & Park, 2019), and more recently, explainability (Shin, 2021), the key concepts of FEAT. FEAT suggests the pressing needs related to socio-cultural, economic, and political concerns with personalized algorithms; however, AA currently lacks a conceptual definition and associated operational measures. While the use of personalized algorithms is prevalent on many platforms, related studies show that most users do not fully understand that platforms, such as *Netflix*, utilize such algorithms to automate recommendations for users (Gran et al., 2021). Hamilton et al. (2014) report that less than 20% of social media users were aware of how algorithms mediate their social media feeds. In other words, the common users are often not aware of how their data are being collected and used and how such personalization algorithms work, let alone pro-actively managing the platforms for privacy concerns. People may be aware of cognate processes without necessarily experiencing the methodological conditions related to algorithms (Koenig, 2020). Research to date on user awareness about the existence and functioning of algorithms is rather limited.

In curating what content is considered personally pertinent, algorithms exert a key role in producing the condition for acceptance, consumption, and engagement in personalized algorithms (Gruber, Hargittai, Karaoglu, & Brombach, 2021). This hidden role points to the need to understand the different levels of AA. The concern regarding how people make sense of algorithms has increased over the last few years (Siles et al., 2020). Eslami et al. (2015) highlight the processes through which users become aware of the functions of algorithms. Awareness can result from active engagement with algorithms on a platform. Over time, active users develop a sense of algorithms through various paths, including the active assessment of algorithms and privacy decisions (Shin, 2021).

Although AA has received research attention, more studies need to be conducted to fully explicate the concept (Gruber et al., 2021). Several definitions for AA have been suggested. For example, Swart (2021) defines AA as the extent to which users are aware of the presence and operation of algorithms in a specific context of consumption and as related to concepts such as fairness, transparency, and trust. Cotter and Reisdorf (2020) define AA as the understanding of what algorithms are, how they are used, how they can benefit people, and how they can negatively impact individuals and certain groups. Zarouali et al. (2021) propose fairness, accountability, and transparency as sub-components of AA.

In addition to these references, explainability has been discussed as yet another component of algorithm literacy (Shin, 2021). A significant challenge for AA researchers is that AI systems are proprietary and remain inaccessible to end-users. Such restrictions make it challenging to establish an objective measurement of awareness. While the objective notion of AA is difficult to establish and seems to differ considerably among populations, it is possible to investigate how people develop AA

and engage in the sensemaking process of AA. As algorithmic processes involve human cognition, behavior, and engagement with the logic of algorithms, users' sensemaking can lead to an algorithmic culture that has a significant impact on how the firms behind platforms with personalized algorithms and users relate to each other (Hargittai, Gruber, Djukaric, Fuchs, & Brombach, 2020).

AA helps users assess and decide how to interact with algorithmic platforms given the basis that educated judgments result in informed decisions (Zarouali et al., 2021). AA also helps users evaluate how platforms, firms, and governments use these technologies, and in doing so, enable them to advocate for responsible technology design to prevent problematic biases and to safeguard user privacy (Akter et al., 2021). AA can enable more users to impact data flows and perceive if or when they and others are being marginalized (Klawitter & Hargittai, 2018). At times, the influence of these efforts may be constrained by a lack of technical knowledge.

In the AI context, user awareness is important because it shapes algorithms as well as user behaviors (Koenig, 2020). Studies have consistently shown that when users are cognizant of algorithms and their functionality, awareness influences how they behave online (Schwartz & Mahnke, 2021). How users think about algorithms and what they know about AI influence the way they interact and engage with algorithms. While some researchers have explored user awareness in various contexts, no standardized scale has been developed yet to measure AA. An important implication from previous studies is that awareness is the result of dynamic engagements with algorithms, meaning that awareness is not static or based on defined knowledge; rather, it is a process of evaluating the algorithm attributes that users experience (Lee, Lee, & Lee-Geiller, 2020). In this light, AA is linked to how well users understand the ethical and normative values of algorithms (Ashok et al., 2022; Hargittai et al., 2020).

Considering that AA includes understanding what algorithms are, knowing where algorithms are deployed, appreciating the intent and goals of those owning or deploying the algorithms, and taking control of user data and privacy (concepts like FEAT) can be considered the factors of AA. Because AA involves critically recognizing the inherent biases and errors in algorithms (Cotter & Reisdorf, 2020), FEAT can be the underlying factors that constitute AA. Per FEAT, AA regards users' understanding of the way algorithms filter and process data, recommends social connections, and reconstructs social realities. Increasing AA is a necessary counterpart to calls for increased FEAT of code as algorithms become embedded in diverse domains of services (Courtois & Timmermans, 2018).

2.2. User-controllable privacy and privacy-preserving algorithms

AI leads to concerns about data privacy and information disclosure as it uses sensitive personal data to train the machine-learning logic—particularly in an era where misinformation is rampant and users' rights are valuable to identity (Lau et al., 2018). To generate recommendations, algorithms should gather personal information, which can be collected by direct requests, possibly triggering privacy risks (Jain, Basu, Dwivedi, & Kaur, 2022). Users make decisions in which they surrender a certain level of privacy in exchange for personalized recommendations that are perceived to be beneficial, convenient, and worthy of the risks.

Privacy is a critical topic regarding algorithms because such platforms require huge volumes of user data (Sundar, Kim, Beth-Oliver, & Molina, 2020). Because data includes users' private information, there are concerns about the privacy implications of AI, such as what is required for user information to be used and what privacy measures are required to protect user data. Privacy in an algorithm context can be defined as the extent to which users are concerned with the potential risks and the right to prevent the disclosure of individual information by firms to others (Lau et al., 2018). This definition renders a right to receive an explanation of outcomes made by algorithms, but this explanatory process of demand leads to a scrutiny of the data used to

train the algorithms (Rader, Cotter, & Cho, 2018). This can result in demands to see the data, which will breach the privacy rights of users from whom the training data were derived (Lau et al., 2018).

Users have the right to know what data are gathered, utilized, analyzed, or otherwise processed and to what extent the personal data are or would be processed (Gutierrez et al., 2019). Ischen, Araujo, Voorveld, van Noort, and Smit (2020) argue that users desire to have some control over how their information is gathered and analyzed, as they have privacy concerns. The importance of empowering people to take control of privacy decisions has increased. Because the interaction with AI is a new feature and due to the black-box nature of AI, users may be more cautious about how their data are taken, used, and represented (Acquisti, Brandimarte, & Loewenstein, 2020).

2.3. PCT

Personalized algorithms rely on user behaviors observed in the data to produce recommendations about what people may like or be interested in Dwivedi et al. (2021). As AI technologies become more pervasive in personalized algorithms, there is a push for the platforms to collect more data to improve the generated suggestions. Inevitably, personalized algorithms pose a serious risk to user privacy. Thus, we use PCT as a theoretical framework for this research to explain users' disclosure of information as a dependent variable.

According to PCT, users weigh the potential benefits and harms in terms of the outcomes in the course of the adoption of new technology. When applying this theory to personalized algorithms, it could be argued that as users weigh perceived benefits (personalizing, accuracy, and relevance) more highly than the uncertainty of risk to privacy, disclosure is likely to occur. The hypotheses based on this theory test the possible effects of risk and the quality of personalization for the decision to disclose information in an algorithm context.

According to Wang, Duong, and Chen (2016), users of platforms with personalized algorithms often face the decision of whether to not share personal information to avert the risk of data mishandling and a violation of their privacy or to share their data to benefit from the personalized recommendations. When users consent to the common terms and conditions of an algorithmic platform, they need to evaluate the benefits and risks of the motivational factors that encourage or discourage information sharing and self-disclosure (Wang et al., 2016).

PCT has been utilized to analyze self-disclosure patterns in various algorithm contexts (Gutierrez et al., 2019). Algorithmic platforms progressively collect information to offer customized services and personalized recommendations. In a recommendation setting, Fast and Jago (2020) argue that when users are not explicitly informed of the strict privacy policy on their personal data, users with a greater orientation of privacy risks are less willing to disclose and share data; however, general users will share personal data when the perceived FEAT exceeds the probable risk (Sundar et al., 2020).

3. Theoretical background and hypotheses development

We propose the key constructs of AA and related concepts based on the Privacy Calculus Theory. The constructs used in this study are contextualized in terms of algorithms. Table 1 describes the contextualized dimensions of these constructs in reference to algorithms and AI.

3.1. Hypothesis development

Algorithm fairness is defined as an approximate equity of negatives or false positives across certain issues or demographic groups (Diakopoulos, 2016). It is increasingly important because as more decisions of greater importance are being made by algorithm-driven platforms and programs, the potential for harm grows. In an attempt to develop fair/equitable algorithms, researchers aim to understand and correct biases, such as by researching the causes of bias in data and algorithms,

Table 1
Summary of the key constructs.

Role	Term	Definitions and operational concepts
AA Components	Fairness	Reasonable and equitable treatment in accordance with accepted rules or principles in algorithms (Shin & Park, 2019).
	Explainability	The availability of explanations about user decision making in relation to algorithms (Rai, 2020; Shin, 2021).
	Accountability	Principle about holding the designers and providers of algorithm-driven platforms liable for the consequences provided by their preset decision-making systems (Diakopoulos, 2016).
	Transparency	The algorithm-generated decisions made by AI should be open, verifiable, and/or visible to the users who use, adopt, and are affected by systems that employ such algorithms (Diakopoulos, 2016).
Antecedent to Efficacy	Algorithm Awareness	The extent to which users are able to understand and correctly assess what algorithms do in a certain AI context as well as their broader impact on how users use and interact with personalized algorithms (Kizilcec, 2016). It includes FEAT (Shin, 2021).
Mediator	Self-Efficacy	Users' ideas on their proficiencies to conduct designated levels of performance and capacities to manage and execute the courses of a task (Hu, Lu, Pan, Gong, & Yang, 2021).
Privacy Calculus	Privacy	Protection intrusion and information access by algorithms in the context of algorithm-driven decision making (Acquisti et al., 2020; Fast & Jago, 2020).
	Trust	Having confidence, faith, and/or hope in algorithms and algorithm-driven decision making (Fast & Jago, 2020).
	Self-disclosure	The act of providing personal information to AI and algorithms (Spanaki et al., 2021; Wang et al., 2016).

defining and applying measurements of fairness, and developing data collection and modeling methodologies aimed at creating fair algorithms. Algorithm fairness is based on its underlying theory of Algorithm Information Processing (Shin, 2021), which posits that algorithm fairness is dependent upon users' perceived algorithm fairness. Relevant research suggests the relation between algorithm fairness and user attitudes, particularly user confidence (Alter, 2021). We propose the causality of algorithm fairness to users' self-efficacy.

H1. Perceived fairness positively influences users' self-efficacy of personalized algorithms.

Explainability is defined as the ability of users to understand how algorithm-driven decisions work (Shin, 2021). A reasonably explainable AI provides the users with easily understandable predictions, increasing the users' confidence in the AI systems. User awareness and an understanding of how and why a particular recommendation is produced and how their user input impacts the result have been found to be significant (Renijith, Sreekumar, & Jathavedan, 2020). Clear transparency and good visibility for relevant feedback improve search performance and user confidence in recommendation systems. The finding of Kizilcec (2016) shows that utilizing explanations can improve positive attitudes and overall satisfaction with a recommendation system. Numerous researchers have reported a causal relationship between explainability and assurance in the context of algorithm-driven platforms and services. Ehsan and Riedl show that human-like rationales promote the feelings of trust, intimacy, rapport, and comfort in non-experts operating AI. Based on existing research, it can be inferred that explainable AI would help users understand the process, thus increasing their faith in the system.

Users are inclined to adopt explainable systems because they can understand how user data are collected and processed and thus how recommendations are generated (Rai, 2020).

H2. Perceived explainability positively influences users' self-efficacy of personalized algorithms.

Algorithm accountability is the notion that the providers of algorithm-driven platforms and services should be held liable for the results of their algorithms (Diakopoulos, 2016). Because algorithms are designed by human developers, personalized algorithms can have built-in human biases and/or mistakes. Accountability for personalized algorithms refers to establishing solid evidence of governance at the organizational level, including well-defined goals and objectives for the system. Relevant research has shown that users develop positive confidence with clearly defined roles, responsibilities, and liability. Algorithm accountability has been theorized along with transparency (Diakopoulos, 2016). The question remains regarding how users perceive algorithm accountability and how it influences users' self-efficacy. Thus, we hypothesize the causality of perceived accountability to the self-efficacy of personalized algorithms.

H3. Perceived accountability positively influences the users' self-efficacy of personalized algorithms.

Algorithm transparency refers to being open with the goal, structure, and/or underlying processes of the algorithms used to search for, recommend, and deliver personalization (Shin, 2021). There have been numerous attempts to examine the effect of algorithm transparency on individuals' psychological attitudes, such as a belief or desire to evaluate algorithms correctly.

Prior research shows that transparency plays a key role in algorithm use by enhancing user trust in an algorithm (Danielsiek, Toma, & Vahrenhold, 2017; Lee et al., 2020). When visible and transparent processes are ensured, users are more likely to consider the results in a more engaging manner (Gran et al., 2021). Open and transparent personalization systems can afford users a sense of assurance, which gives the user a sense of efficacy. Given this, H4 is proposed:

H4. Perceived transparency positively influences the efficacy of personalized algorithms.

User efficacy is defined as the degree of one's feelings about his/her ability to accomplish goals (Bandura, 1986). Research has shown that efficacy plays a key role in how humans perform because it directly influences factors such as motivations and goals, affective tendencies, perceptions of opportunities, and outcome expectations in social environments. In the context of algorithms, extensive research has shown that user efficacy is significantly associated with user adoption and the experience with machine-learning systems (e.g., Mullins & Cronan, 2021). In particular, the association of efficacy and privacy has been found to be significant (Chen & Chen, 2015). As machine-learning algorithms increasingly rely on data to generate personalized decisions, algorithm efficacy is considered a pivotal factor that is related to willingness and the ability to perform a privacy evaluation. Therefore, it is worthwhile to test the effect of users' self-efficacy on the privacy of personalized algorithms. Given this, we propose H5:

H5. Efficacy positively influences the perceived privacy of personalized algorithms.

As efficacy is understood as a factor that represents an individual's perception of his/her capacity to make algorithm related decisions, there is the possibility of a relationship with self-disclosure. With high efficacy, people would be more confident in their decisions regarding what and how to share data with algorithms. The influence of efficacy on self-disclosure can be best conceptualized when one gains a complete understanding of how the algorithmic systems work and how to work with the systems. It has been debated that algorithm efficacy should be understood from a user perspective in a specific context of disclosure

behavior. User's efficacy can play a key role as a bridging variable to help show the connection between self-disclosure and personalization in algorithmic systems (Danielsiek et al., 2017). With an understanding of efficacy, one could understand how users who have high efficacy with a strong understanding could possibly have a high capacity of self-disclosure in algorithmic environments. Efficacy positively influences users' self-disclosure on social media (Chen & Chen, 2015). It can be proposed that users' sense of efficacy leads to information sharing and self-disclosure to platforms with personalized algorithms. Therefore:

H6. Efficacy positively influences users' self-disclosure with personalized algorithms.

The effect of efficacy on trust has been widely touted and researched in a variety of contexts (Reisdorf & Blank, 2020). Trust is a concept related to an individual's desire to have high confidence in algorithms. Trust occurs when a person believes in the reliability of the trusted systems. When users have high efficacy and thus know the functionality and the inner workings of the algorithms, they tend to trust the AI. Establishing trust is a dynamic process involving cognitive efficacy based on a user's confidence and competence in algorithms (Lee et al., 2020). For example, when people are to make a decision related to algorithms, they would prefer decisions based on the choices of algorithms they trust more. It has been found that efficacy is highly related to trust (Danielsiek et al., 2017). In particular, self-efficacy significantly influences trust in personalized algorithms (Hasan et al., 2019), as the more users know about algorithms, the more trust they develop (Hu et al., 2021). Therefore, H7 is advanced:

H7. Efficacy positively influences users' trust in personalized algorithms.

Privacy is defined as the state of being free from intrusion or disturbance in one's private affairs when using algorithmic systems (Fast & Jago, 2020). Content personalization based on algorithms affords a series of values to users, but they entail a huge amount of user data in return. Privacy can be compromised because algorithms can be misused or abused. Thus, in algorithm contexts, trust, privacy, and information disclosure are closely related concepts (Fast & Jago, 2020). Privacy is considered a user right, and as such, the disclosure of personal information is controlled. Perceptions of privacy in algorithms are significantly conditioned by factors in the algorithmic environment, particularly factors such as personalization and trust (Lau et al., 2018). Recent studies have suggested a link between privacy and self-disclosure, and thus we propose H8:

H8. Privacy positively influences users' self-disclosure.

In the algorithmic recommendation context, trust is defined as dependable confidence in the legitimacy of algorithm-driven decision making and for user readiness to accept the algorithmic system's competencies. A user's intention to disclose personal information hinges upon a privacy security analysis (Hasan et al., 2019). Per PCT, users balance perceived privacy and expected values (Gutierrez et al., 2019). Users are increasingly paying more attention to their data and becoming uncomfortable with how data about their preferences, search histories, locations, and interests are used by personalized algorithms (Sundar et al., 2020). High levels of trust in algorithms can alleviate users' privacy concerns and information disclosure. Trustworthy recommendations afford users a sense of privacy, which in turn encourages the willingness to disclose personal information. Relevant research has found that cognitive processes, such as users' assessment of normative values and trust, are determining factors of users' decisions to disclose personal information (Ashok et al., 2022). Based on the literature, we advance H9 (Fig. 1):

H9. Trust positively influences users' self-disclosure.

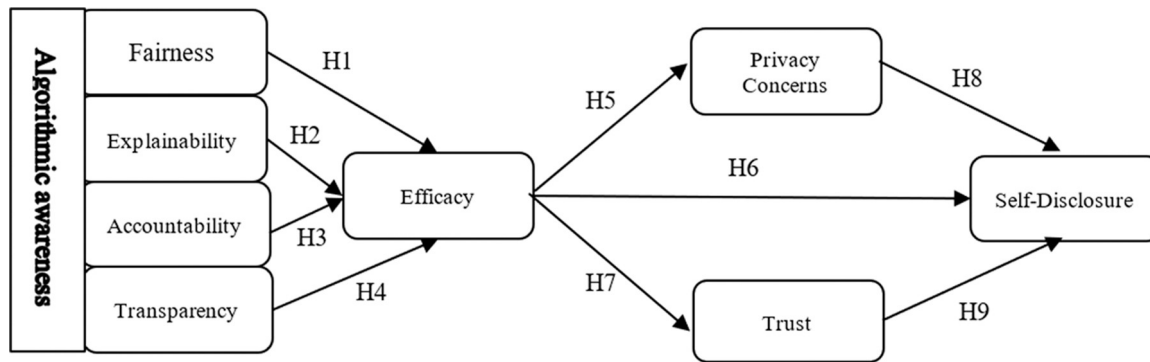


Fig. 1. The Privacy Calculus process for algorithms (Source: Adapted from Shin & Park, 2019).

4. Methodology

The partial least square method of structural equation modeling (PLS-SEM) was used to explain the structure model and to examine the interrelationship among factors.

4.1. AA scale development

We developed composite scales to measure the construct of AA based on the a priori set of related measurements based on a systematic literature review. We developed and confirmed the AA-instrument with composite scales based on the three steps proposed by Zarouali et al. (2021) in which we (1) classified a pool of possible elements associated with the sub-components of algorithm knowledge and understanding; (2) drafted the AA-instrument to evaluate the content and face validity of these elements, and (3) confirmed the instrument using a nationally representative sample in South Korea. In this last step of the pre-test, we allocated the whole sample into three sub-groups, representing three different news aggregators: *Naver*, *Kakao*, and *Daum*. Each sub-group conducted a specific task: an exploratory goal, a confirmatory goal, and a generalization goal. Fig. 2 illustrates the steps involved in developing and validating AA.

4.2. Questionnaire development

Based on the hypothesis of the measurement sub-model, the scale indices previously applied, and relevant theories, a questionnaire was developed to measure AA and its related factors (see Appendix for measurement items). The measurements in the model were developed based on previously validated instruments. The FEAT measurements were modified from Renijith et al. (2020), Rai (2020), and Shin et al. (2020). The measurements of self-efficacy and self-disclosure were based on Danielsiek et al. (2017) and Sundar et al. (2020), respectively. The trust and privacy items were modified from Rai (2020) and Lau et al. (2018) respectively. As SEM was used to validate the model, a seven-point Likert scale was appropriate to use for respondents to indicate their degrees of agreement with the items. Before the formal survey, the questionnaire was pilot-tested, and 45 pre-survey responses were received. The questionnaire was then revised and finalized based on the pilot data. Wordings were modified to improve clarity, and one item was dropped. The final questionnaire comprises 24 items grouped under three dimensions: awareness (12 items), efficacy (3 items), and privacy (9 items).

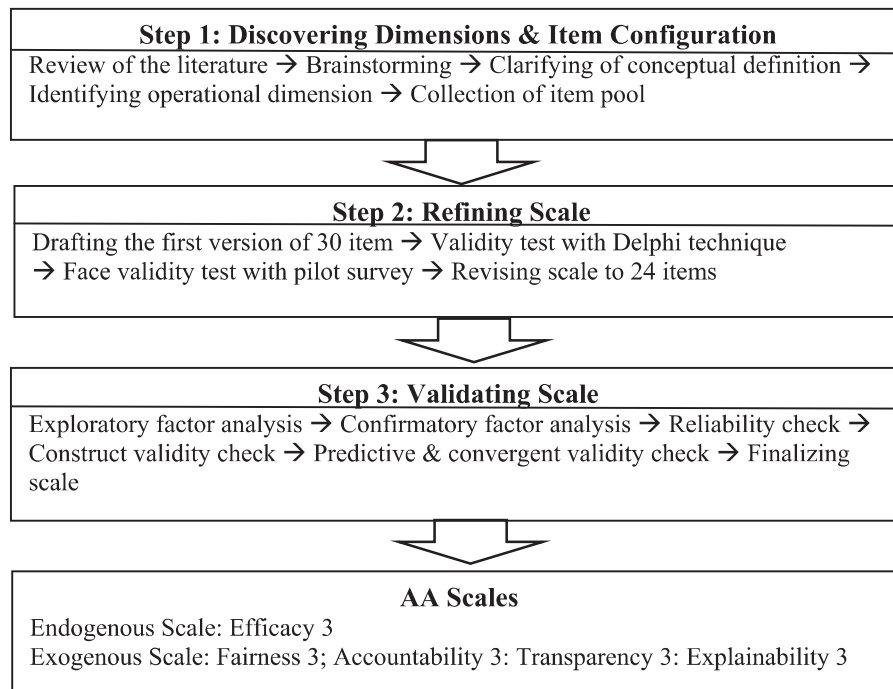


Fig. 2. Development and validation of the AA-instrument (Source: Adapted from Zarouali et al., 2021).

4.3. Reliability and validity analysis

The Cronbach's alpha scores revealed that all the composite scales were reliable (scores between 0.80 and 0.90), which means there is sufficient internal consistency (Table 2). A confirmatory factor analysis was conducted to check the results of the exploratory factor analysis. The analysis showed that the factor loadings for all eight composite scales were consistently high and positive (above 0.783), providing evidence of convergent validity and discriminant validity. A series of Pearson correlation tests were conducted to check for reciprocal relationships among the eight factors. A variation test for multicollinearity showed no sign of collinearity. The square root of the average variance extracted (AVE) from the latent variable in all cases exceeded the correlations between each pair of factors, indicating discriminant validity. In all cases, the AVE was greater than 0.70, and all items had acceptable convergent validity because the results showed that all parameters were statistically significant.

4.4. Data collection

PLS-SEM has been widely used due to its ability to examine non-normal data distributions (Hair, Sarstedt, Hopkins, & Kuppelwieser, 2014). PLS-SEM provides excellent tools to perform a simultaneous test for complex associations among the variables in the case of heterogeneous and multivariate phenomena (Bollen, 1989). Based on the deployment of the PLS Approach, we used the SmartPLS (version.3.3.3) software tools for PLS-SEM. While the PLS method can work with a small sample size, it is suggested that a sample size between 100 and 400 is acceptable when complex structure modeling is undertaken (Hair et al., 2014).

A national survey was conducted during the fourth quarter of 2021 with users of news platforms with personalized algorithms who had experience of at least six months prior to completing the survey. This was designed to ensure that respondents had a certain level of experience with, exposure to, awareness of, and attitudes toward personalized algorithms. In measuring users' AA, we assessed their perceptions and self-reported knowledge of algorithms by inquiring about their implicit/explicit understanding of algorithms: explicit algorithm usage time, functional and technical understanding (computable functions), and implicit FEAT issues, such as an appreciation of the way algorithms filter

and process data, generate recommendations, and create personalized predictive content for users. The pilot and pre-survey questionnaires were assessed by four professors and experts who offered concrete suggestions. The respondents were informed of FEAT in the specific context of personalized news algorithms because these concepts are vague and may possibly be outside of common definitions for laymen who do not specialize in algorithm research or practice. A total of 450 questionnaires were disseminated, and 397 valid responses were collected from a large public university in South Korea (Table 3). The sample was based on enrolled students in classes on topics such as digital media, algorithms, and user design. Screening interviews were conducted to measure the degree of personalized algorithm usage. The preliminary survey included four questions: (1) Which news platform are you currently using? (2) How much time per week do you spend on the news platform? (3) What do you use the news platforms for? and (4) What content did the news platforms recommend to you? We used major news aggregator platforms (*Naver*, *Kakao*, and *Daum*) currently used in Korea based on the market share report by the Korean Communication Association.

4.5. Procedures

In the experiment, respondents were required to recall their experiences with personalized news aggregators, such as *Naver News*, *Kakao News*, and *Daum News*, as an example of a platform with personalized algorithms. The respondents were given the opportunity to use and search via news platforms, submitted their data on user characteristics, preferences, demographics, and default data on certain items, and then evaluated the extent to which the contents were personalized/relevant to their profiles. They also were allowed to search their preferred online platforms for shopping items, such as books, drama, films, video clips,

Table 3
Descriptive statistics.

Characteristics	Sample
Age (Mean/SD/Median)	29.48/5.87/34
Gender (Female Rate)	50.19
College Educated (%)	30.11
Algorithm Platform Experience	2.1 years

Table 2
Confirmatory factor analysis.

Factors	Mean	Standard deviation	Cronbach's alpha	AVE	Composite reliability	Factor loading
Transparency	4.61	1.040	0.851	0.773	0.811	0.813
	4.64	1.117				0.741
	4.43	1.070				0.903
Accountability	4.35	1.187	0.843	0.734	0.734	0.841
	4.61	1.211				0.851
	4.06	0.979				0.901
Fairness	4.04	1.274	0.820	0.724	0.714	0.731
	4.44	1.223				0.913
	4.07	1.084				0.844
Explainability	4.04	1.223	0.868	0.711	0.743	0.831
	4.14	1.272				0.831
	4.16	1.264				0.911
Efficacy	4.36	1.118	0.906	0.703	0.841	0.814
	4.23	1.058				0.742
	4.13	1.048				0.853
Privacy	4.17	1.009	0.806	0.712	0.801	0.945
	3.82	1.079				0.901
	4.08	1.062				0.821
Self-Disclosure	4.40	1.129	0.871	0.704	0.704	0.734
	4.10	1.282				0.823
	4.19	1.229				0.822
Trust	4.15	1.130	0.829	0.832	0.921	0.835
	4.33	1.147				0.745
	4.51	1.203				0.911

and music.

4.6. Data validation

Table 4 shows the means, Pearson correlation coefficients, and standard deviations among the variables. AA was positively correlated with FEAT: fairness ($r = 0.159$, $p < 0.01$), explainability ($r = 0.301$, $p < 0.01$), accountability ($r = 0.16$, $p < 0.05$), and transparency ($r = 0.49$, $p < 0.01$). Furthermore, efficacy was positively correlated with privacy ($r = 0.39$, $p < 0.01$), trust ($r = 0.29$, $p < 0.01$), and self-disclosure ($r = 0.28$, $p < 0.01$).

The results for the model's goodness of fit indices are within the recommended ranges (Table 5). Selected goodness-of-fit indices were evaluated with suggested cutoff scores. The Variance Inflation Factor value of the SEM in this study was between 1 and 3.141 (less than 5), indicating no sign of collinearity among the study constructs (Hair et al., 2014).

5. Results

In estimating the structural relations in the model, we used the estimated coefficient. The PLS analysis revealed that all the nine paths were meaningful, as all the coefficients and t -values were significant (Table 6 & Fig. 3). FEAT has a significant effect on users' self-efficacy (β : 0.256, 0.279, 0.277, 0.412; t -value: 3.755, 2.630, 5.571, 5.940) and directly influences privacy (β : 0.681, t -value: 12.44), trust (β : 0.675, t -value: 10.74), and self-disclosure (β : 0.698, t -value: 12.399). About 59% of the variance in efficacy was accounted for by FEAT. The effects of privacy-trust on self-disclosure were significant ($\beta = 0.233$; t -value: 5.052).

About 44% of the variance in self-disclosure is explained by trust, privacy, and efficacy. The explanatory effect value f^2 of FEAT to efficacy is 0.480, which shows large-effect explanatory abilities (Table 7). The explanatory effect value f^2 of trust and privacy to self-disclosure is 25%, which shows a large effect explanatory power. The explanatory effect of privacy also shows a large effect. Therefore, the exogenous variables are very good at explaining the endogenous variables with a high degree of explanatory effect value.

The mediation effects of privacy-trust on the path from efficacy to self-disclosure were tested using Preacher and Hayes (2008) bootstrapping methodology. A bootstrapping analysis with 5000 resamples was performed to test the mediating effects. This method produces the confidence intervals (CI) of indirect effects. There is no need for normal distributions (Table 8). The result was statistically significant, which also supports the mediation effect of privacy-trust on the liaison between efficacy and self-disclosure. The total effect ($\beta = 0.635$), which was computed by conjoining the direct effect ($\beta = 0.620$) with the indirect effect ($\beta = 0.015$), was statistically significant. The test results indicate that privacy-trust mediated the significant positive relation between efficacy and self-disclosure, supporting the facilitating role of AA.

Based on the results, it can be concluded that the effect of efficacy on self-disclosure is mediated via privacy and trust. The role of the privacy-trust factor as a mediator may help platforms in identifying and utilizing

Table 5

Model fit indices.

Fit statistics	Structural model	Results	Recommended value
χ^2/df	1120/ 237 = 4.72	Acceptable	< 5 (Hair et al., 2014)
SRMR (RMS theta)	0.05	Acceptable	< 0.08 (Hair et al., 2014)
RMSEA	0.060	Acceptable	0.05 < α < 0.10
CFI	0.920	Acceptable	> 0.90 (Bentler, 1990)
NFI	0.931	Acceptable	> 0.90 (Bentler, 1990)
IFI	0.900	Acceptable	> 0.90 (Bollen, 1989)
TLI	0.920	Acceptable	> 0.90 (Hu & Bentler, 1999)
RFI	0.894	Acceptable	> 0.80 (Hair et al., 2014)
AIC	659	Acceptable	The lower the better (Hair et al., 2014)
Hoelter's Critical N	227	Acceptable	> 200 (Hoelter, 1983)

Table 6

Results of the path analysis.

H	Coefficient	S.E.	t -value	P	Result
H1	0.256	0.029	3.755	0.00***	Supported
H2	0.279	0.070	2.630	0.009**	Supported
H3	0.277	0.057	5.571	0.00***	Supported
H4	0.412	0.070	5.940	0.00***	Supported
H5	0.698	0.054	12.399	0.00***	Supported
H6	0.499	0.098	5.815	0.00***	Supported
H7	0.850	0.064	8.592	0.00***	Supported
H8	0.233	0.055	5.052	0.00***	Supported
H9	0.348	0.150	4.083	0.00***	Supported

* $p < 0.05$; ** $p < 0.01$; *** $p < 0.001$.

factors that may determine algorithm acceptance and use. The role also suggests the use of the privacy calculus process in algorithm contexts.

6. Discussion: empowering users to control their privacy and data

The findings of the study provide proof-of-concept insights for developing the AA processing model in an algorithm context. The model clarifies that interacting with algorithms involves AA processes wherein algorithm attributes are used to formulate a heuristic of user motivation and to trigger actions in algorithm adoption behavior. The findings are consonant with previous assumptions of personalized algorithms: (1) users' lack of underlying knowledge of how algorithms work (Hargitta et al., 2020); (2) the black-box nature of algorithms within everyday AI consumption (Eslami et al., 2015); and (3) limited knowledge on algorithms impeding users in making informed privacy decisions (Shin, 2021; Swart, 2021). Our AA model not only confirms these existing findings but also extends them by clarifying how users make sense of algorithm attributes, how users establish a sense of efficacy, and how the acquired awareness influences privacy and leads to self-disclosure. The findings of this study offer meaningful insights into the relationships

Table 4

Correlation of constructs.

Variable	Fai.	Exp.	Acc.	Tran.	Eff.	Tru.	Pri.	Dis.
Fairness	0.86							
Explainability	0.02	0.72						
Accountability	0.10	0.14	0.90					
Transparency	0.11	0.10	0.09	0.82				
Efficacy	0.14*	0.29**	0.38**	0.21*	0.79			
Trust	0.05	0.04	0.10	0.08	0.35**	0.89		
Privacy	0.01	0.08	0.26*	0.15*	0.01	0.37**	0.90	
Disclosure	0.07	0.15	0.08	0.20	0.35**	0.31*	0.20*	0.88

Note: Diagonal elements are the square root of AVEs for each construct. * $p < 0.05$; ** $p < 0.01$.

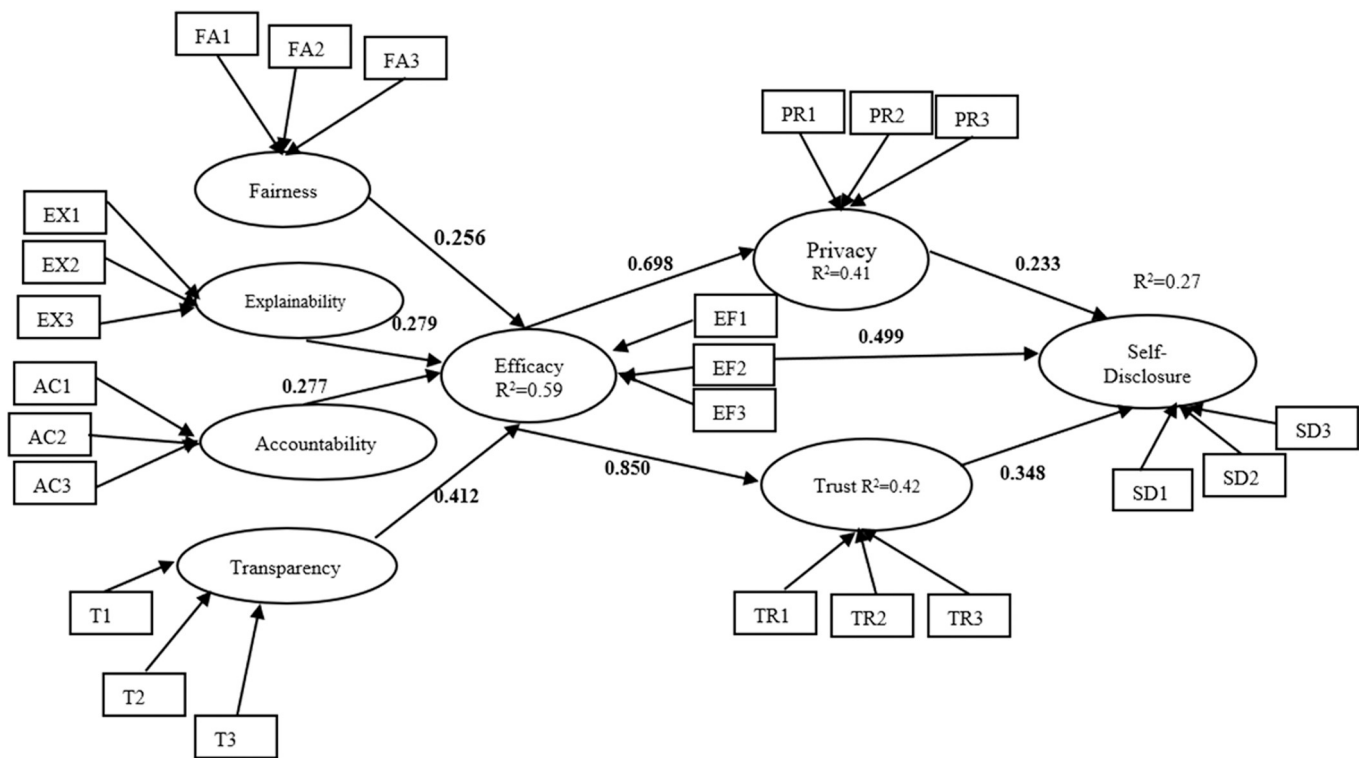


Fig. 3. The results of the algorithm Privacy Calculus process.

Table 7
Value of R^2 and f^2 .

	R^2	R^2 adjusted	f^2	Effect size
Efficacy	0.682	0.599	1.322	Large
Privacy	0.487	0.410	1.309	Large
Trust	0.472	0.424	0.187	Large
Self-Disclosure	0.278	0.250	2.342	Large

Table 8
The results of simple mediation.

	Estimate	S.E.	Lower CI	Upper CI	Indirect effect %
Privacy	-0.0628	.0182	-0.1245	-0.0139*	15
Efficacy	-0.1492	.0374	-0.2344	-0.0831	23
Information Disclosure	-0.139	.0474	.2843	.1103	15
Trust	-0.0517	.0125	-0.1174	-0.0128*	28
Efficacy	-0.1492	.0345	-0.2244	-0.0729	12
Information Disclosure	-0.133	.0259	.3582	.1049	40

* Significant mediator.

among awareness, trust, and privacy in algorithmic platforms. The discussions can be analyzed in tandem with existing literature in the following aspects.

First, awareness related to technology has been researched as a critical theoretical portion of various digital literacies. Advancing from existing literature (e.g., [Cotter & Reisdorf, 2020](#)), this study identified the types of AA practices users engaged in with algorithms. While previous works have proposed the importance of AA and/or have identified the gaps in understanding algorithms, we identified FEAT as a characteristic of AA and clarified a heuristic for AA in the use of algorithmic platforms. The AA scale can lead to a theoretical refinement in the studies of behaviors and perceptions associated with algorithms. AA is

proven to be a key antecedent of efficacy as it is expected to predict privacy and cause positive disclosure decisions. AA predicts users' trust perceptions toward platform algorithms. Our conceptualization of AA in reference to FEAT gives important implications that algorithm efficacy goes beyond "know-what" and is instead the pursuit of recognizing the context, meaningfully controlling AI, and acting to meet the justifications of responsible AI. Furthermore, AA can serve as a key theoretical concept in its terms as many related concepts can be derived from it, such as algorithm literacy, algorithm aversion, and algorithm adherence ([Zarouali et al., 2021](#)). Having awareness of algorithms can be a key initial step to ensuring that users can make informed privacy evaluations and decisions ([Ahmad, Widén, & Huvila, 2020](#)). One's ability to recognize and understand the potential of algorithms/AI for success is critical.

Second, the findings clarified the specific roles and processes of FEAT in users' behaviors related to algorithms. The findings show that AA influenced users' trust and attitude through two different routes of cognitive processing: heuristic and systematic processing. This is accomplished through FEAT heuristics first and through a systematic process based on trust second. Heuristic processing involves simplifying the assessment of FEAT to quickly evaluate the service quality. Systematic processing involves thoughtful processing of personalization and accuracy. A heuristic process is less resource demanding and less analytical as users normally do not have the expertise to evaluate specialized algorithm features, whereas a systematic process requires more effort and is more deliberate (whether recommendations results are accurate, predictable, and customized) based on the established trust. The models reveal that FEAT is used as a heuristic tool to assess trust in an algorithm. Users' heuristic process of AA influences their trust, and increased trust influences the systematic processing of performance expectancy, which is positively associated with attitude and intention ([Sundar et al., 2020](#)). Just as the issues of FEAT are considered essential values in social systems, they are in algorithm-based platforms as well. The qualities of FEAT not only play a key role in stewarding trust but in guiding users' evaluations of performance in terms of the usefulness of the algorithmic predictions ([Monzer et al., 2020](#)). Algorithm

users develop their own cognitive processes of AA based on FEAT. User reactions to perceived performance depend on or at least closely related to how users perceive, understand, and process information regarding FEAT as AA. This relation can be explained as heuristic insofar as users rely on their perceptions of FEAT to determine their feelings of personalization in algorithmic platforms. That is, users determine performance value in the context of an algorithm according to their perceived FEAT of content.

Third, by applying PCT to personalized algorithms, we showed that the construction of users' self-efficacy is a psychological antecedent to the privacy calculus process by which users seek to assess privacy concerns and the decision to self-disclose. The roles of efficacy in self-disclosure and trust have remained unexplored and largely unknown, particularly in algorithm contexts. Very few studies have examined the psychological antecedents of privacy calculus, and even these studies have largely overlooked the cognitive process of efficacy by considering users' self-efficacy as a product (Culnan & Armstrong, 1999; Danielsiek et al., 2017). We approached algorithm efficacy as a process of perception that guides users' heuristic and systematic efforts to learn and to evaluate the privacy and establish trust. The findings suggest that algorithm efficacy is the user's ability to knowledgeably decide on self-disclosure of personal information, and from there, effectively establish trust in personalized algorithms. Thus, algorithm efficacy is not a singular activity; rather, it involves multiple paths and shapes (Zarouali et al., 2021). The model shows that the level of efficacy influences users' privacy and trust through users' privacy calculus processes. The finding highlights the significance of efficacy as a facilitating mechanism, illustrating relations between AA and subsequent attitudes toward disclosure decisions. This relation illustrates that knowing is one thing and accepting algorithms is another. Users' awareness should be translated into efficacy, which then triggers subsequent actions of privacy evaluation and informed self-disclosure. This argument is supported by the mediating role of privacy in the relationship of trust to self-disclosure and constitutes theoretical contributions to PCT through the conceptual refinement of the basis of privacy calculus (Sundar et al., 2020), how AA is formed (Shin, 2021), how it is mediated by privacy (Ischen et al., 2020), and how self-disclosure is influenced by users' self-efficacy (Hamilton et al., 2014). By focusing on the unique aspects of user decision-making process, we illustrated the algorithm sensemaking under which different types of processes are evoked, how they interact (H1, H2, H3, H4) and how they apply to users' processing of trust (H7), the evaluation of privacy (H5), and self-disclosure behavior (H6 & H8). We have elaborated how algorithms shape not only the nature of sensemaking but also the form of sensemaking processes.

Fourth, the findings clarified the dimension of trust in AA. The identified role of trust in an algorithmic platform implies that people do trust the recommendations of an algorithm when their awareness test is complete with FEAT. From the algorithm information processing in the model, it can be inferred that AA and performance values are positively associated with trust (Reisdorf & Blank, 2020). The mediating effect implies that efficacy can lead to information disclosure only through an experience with privacy and trust. Trust connects the dual heuristic and systematic processes, linking the heuristic and systematic mechanism (Shin & Park, 2019; Siles et al., 2020). This trust link can be a key clue to understanding algorithm qualities, algorithm experiences, and users' interactions with AI. Certain algorithm features provide users with clues for trust, and trust allows users to adopt algorithms with feelings of usefulness and efficacy (Gutierrez et al., 2019; Lau et al., 2018). Trust formed through heuristic processing is more likely to have cognitive attributes that reflect the AA valuation, whereas trust shaped through systematic processing is more likely to exert effects on performance evaluations due to a lack of understanding of FEAT. Findings on the role/process of trust and the relationships among its related measures not only corroborate the theory's core argument—that cognitive decision is guided through a two-track process—but also advance the theory by linking this process to the two-step flow of sensemaking—the dual

process by which users create an understanding so that they can act in an informed and principled way. Users actively process and proactively control stimuli and algorithmic curations, evaluating privacy based on normative values and trust.

Lastly, the findings highlight the active role of users in interacting with algorithms. Our conceptualization of AA from an algorithm sensemaking perspective is based on the notion that users are active agents engaged in their own cognitive development and can have an influence on algorithms through their informed actions (Min, 2019). Key to this notion of an active agent is the belief that individuals acquire self-efficacy, which enables them to exercise a measure of control over their feelings, perceptions, and actions, and that what people think, believe, and feel affects how they behave (Swart, 2021). Our findings support a proposition of user behavior in which the beliefs that people have about themselves are critical rudiments in the design and exercise of algorithmic platforms (Shin, 2021). Thus, users are considered both the products and producers of their own algorithms and of their AI systems. Humans are consumers as well as producers of AIs as the algorithms show what users want to see and what is relevant based on users' own sensemaking results. This conceptualization makes the sensemaking theory applicable to human interaction with algorithms.

6.1. Theoretical contributions and implications

Our work contributes to the ongoing discussion on awareness, trust, and privacy concerns in the context of algorithms. We newly conceptualized AA along with FEAT and tested the heuristic dimension of AA. Our results contribute to theoretical knowledge by clarifying what constitutes AA, how AA works, and what affects AA during the use of algorithms as well as how trust can be established, enhanced, and measured. The results add to the theoretical understanding in the following ways.

First, our results made a theoretical refinement in privacy by clarifying how users weigh privacy and personalization in the context of algorithms (H5 & H8). To the best of our knowledge, this is the first study to link privacy concerns with user awareness and information disclosure. The previous theory of privacy concerns largely downplays contextual factors by focusing on the mechanical calculation of risk-benefit (Gutierrez et al., 2019). Our novel findings suggest the intricacies of privacy concerns, revealing that various situational factors should be examined to better understand users' decision making regarding self-disclosure to algorithms. User awareness provides a context for understanding the capacities of algorithms, humans, and their co-evolving and inter-defining relationship (Shin, 2021). The contextual factors include the way in which information is collected by the platforms, how accountable and transparent the algorithms are, and the trust perceptions formed during an interaction with a specific algorithm, among others (H1, H2, H3, H4). Our new findings suggest that the effect of algorithm privacy concerns is very likely to be outweighed by these contextual factors at a specific level, i.e., related to a specific platform provider in terms of FEAT. AA should include contextual issues and other inputs to understand the meaning and significance of data and behaviors related to algorithm: users' interpretations of the path algorithms translate meaning, influence us, structure our interactions with AIs and the processes affecting what we consume, and affect how we derive meaning from algorithms and what we think.

Second, the procedural dimensions of AA lend themselves to users' privacy calculus processes in algorithms. The algorithm sensemaking perspective differs from previous theories in that the information processing model overemphasizes the influence of mechanical factors in algorithm design and development (Min, 2019). While traditional theories acknowledge the effect of evolutionary aspects in users' roles and variability, they downplay the role of constructivism, which considers algorithm decisions as the outcome of evolved algorithmizing (Swart, 2021). Algorithm privacy calculus as defined in this study espouses a bidirectional influence or co-evolving nature in which interactions

adjust human development such that users become able to engage with increasingly complex and imperceptible algorithmic processes. In turn, the algorithmic processes establish a new selection mechanism for the generation of personalized algorithmic systems. This bidirectional influence results in important contributions to algorithm privacy calculus processing theory.

6.2. Implications for practice: user awareness by design

One of the important practical contributions of this study is that it highlights the importance of user awareness in algorithm design and operation. As AI is being rapidly developed, the industry must develop ways to design algorithms to be human-aware and user-controllable. As AI continues to impact the mode in which we interact with technologies, how to design AI in a way that allows users to have meaningful control over algorithms, how to warrant transparent interaction and fair algorithms, and how to embed FEAT in the interface will be pressing questions to address. There is a growing need for informing AA, and those who develop algorithms should be trained in ethics and required to design a code of ethics that considers societal processes and their interactions with context. Algorithmic platforms should have a strategy for how to communicate and how functioning is transparent, fair, impartial, and in line with accepted social norms among users. The industry can use the identified AA components as matrixes or criteria for analyzing user activity and behavior for platform design purposes. The AA components can serve as a baseline to understand users' readiness and how they are empowered to find a right balance between personalization and privacy when using algorithmic platforms.

Our AA framework offers practical guidelines for informed AA practices that could be integrated into personalized platforms. Issues of FEAT have been essential issues in AI, and users seek assurances on such issues in using AI. A meaningful implication of this is that the industry can design innovative "users-in-the-loop" algorithmic platforms to leverage users' capability to deal with algorithm problems and limitations. The AA matrix can improve algorithm performance objectives by engaging human users in reflecting their heuristics, capabilities, and preferences. Developing user-centered algorithmic services requires the integration of users' sensemaking processes along with the capability to reflect these processes in an algorithm's design.

As the findings suggest, the functional features of algorithms are processed through users' sensemaking and processing regarding users' self-efficacy, which is mediated by trust. Awareness thus facilitates the cognitive evaluation of norm, performance, attitude, and intention. Our findings suggest that awareness encompasses a system with knowledge of privacy in which it is aware of itself as well as the contextual situation. This result implicates "artificial consciousness" as a new area of AI. The concept of artificial consciousness can be concretized by incorporating AA. As our results of AA provide a point of reference for artificial consciousness, future studies may further investigate this consciousness with the theoretical justification of FEAT as a construct of AA. Understanding the role of consciousness and AA would give insights into the development of a user-centered interface for algorithmic systems.

Our results related to privacy also provide operational insights into how to utilize the user privacy calculus process in algorithms. There has been increasing interest in the problem of building accurate personalization over aggregate data while protecting privacy at the level of individual data. Our findings of the user privacy calculus process provide implications for designing privacy-preserving algorithms. In personalized recommendations, for which the goal is to expand the connections between users and increase their engagement with algorithms, it is essential to allow users to engage in user-driven and transparent privacy calculus. Drivers of self-disclosure should be examined from a procedural view, where users make awareness evaluations to decide whether to share their personal information to obtain the benefits they will derive from personalized recommendations. The industry can benefit from utilizing our privacy calculus process in designing privacy-

preserving algorithms, which can enhance user experience and increase users' engagement in algorithms while guaranteeing privacy for users.

Regarding trust, algorithm designers can seek guidelines for how to implement better trust strategies to attract more consumers. For instance, based on our results, the industry can build a better reputation for their services and systems through a transparent process, conveying a sense of transparency and trust, implementing privacy policies to protect individuals' information privacy, making efforts related to service comfort through efficacy, providing higher quality recommendations that can satisfy users' demands, supplying appropriate measures when privacy intrusion occurs, etc. Algorithm marketers can also be guided to promote platform services adoption. For example, they can devise more accountable strategies in that algorithm firms collect and use individuals' personal data. In addition, they can make efforts to assist users in understanding algorithmic systems and provide effective incentives for them to adopt the services.

In terms of user awareness, we have clarified how users make sense of privacy concerns and what types of AA practices users engage in with algorithms. When users believe in the function of personalized algorithms, warrant their data, and have strong confidence about privacy protection, users tend to share more information. The results can be translated into practical ends in that when users' awareness increases, their willingness to disclose information increases. This result provides practical implications that user awareness is based on FEAT and further indicates how it triggers and facilitates the privacy calculus process. The results give insights into the awareness that AA is more than a simple realization of algorithm, its presence, and its immediate consequences of personalization and predictions. AA is more related to the fundamental underlying issues of normative values and less to the performance or results of algorithms. The industry can use FEAT as a heuristic cue for AA in the assessment of platforms with personalized algorithms.

For user-controllable AI, to allow users more meaningful control over algorithms, the industry can develop AA strategies based on the two practical principles: (1) with a subdimension of FEAT, AA can be considered in terms of procedural aspects rather than as a performative consequence of algorithms. AA can then be said to be users' perceived notions about normative values and practices of algorithms rather than material terms of algorithm performance or technicality; (2) AA can be said to be a process rather than a product or one-time finality. Awareness as a process refers to users engaging in activities that involve becoming conscious of algorithmic processes. Rather than considering it from a product view, AA serves as a privacy calculus process of evaluation in which perception, heuristics, systematic evaluation, and action-taking happen concurrently. AA can be a statement of a user's mind and how users perceive algorithms and thus how they develop a sense of efficacy. In this light, AA is best understood and put into use as a set of social practices in terms of the ways users experience algorithms in their everyday lives along with the interactions mediated by actual algorithmic services. Algorithmic platforms can be considered experience systems in which their use enables users to learn how a specific algorithm works. Users' perceptions and their psychological state of mind are critical in rationalizing how and why users perceive and feel the way they do about issues related to AI as well as how they accept and experience AI services.

6.3. Limitations and future research directions

Although innovative and novel, our results are not conclusive and suggest that the measures of awareness, FEAT, and privacy are not yet comprehensive enough for the empirical study of AA. While the relationships identified through the model serve as a first step in achieving long-term goals of meaningful user-controlled AI, such relationships need to be confirmed in various contexts of AI. As one of the first data-driven modeling studies, this study relied on a preliminary conceptualization and a basic operationalization of AA informed by the known,

basic factors that influence algorithmic curations. While the AA dimensions that are included in the instrument and composite scales are based on an extensive literature review and relevant empirical work, we do not claim that this is an exhaustive list of dimensions. There may be other dimensions that are part of the awareness construct that are not included in the instrument. Future empirical studies may uncover relevant components that can be added to the instrument. Future research could probe the theoretical and methodological underpinnings of FEAT and awareness in greater detail. Future studies could adapt our awareness framework to further examine algorithm behaviors and meaningful user control in various domains. The extent to which user awareness is influenced by or influences algorithm values/behaviors would be interesting topics to pursue.

7. Conclusions

The results show that AA leads users to envisage, perceive, and interact with algorithms, depending on their understanding of the control of the information flow embedded within them. The awareness that users have of algorithms influences the trust of algorithmic processes and the way users evaluate privacy concerns and their self-disclosure. Only when users are fully cognizant of algorithm practices can they assess privacy, determine information sharing, and appreciate how algorithms affect the data they receive from users. Users' sensemaking processes of AA provide conceptual frameworks for algorithm system design and a practical guideline for the design of personalized algorithms.

Algorithmic biases and the issue related to the opaque processes of AI put forward meaningful user control as one of the key principles concerning responsible and fair AI. We contribute to this scholarly debate by identifying a set of four properties aimed to support AA practice to design and develop systems that can incorporate a meaningful form of human control. In this article, we have shown how users make sense of the workings of personalized algorithms, how they come to acquire awareness, and what the implications of their efficacy through

awareness are for addressing privacy concerns. A critical way for users to gain more influence over algorithmic processes is by becoming aware of them. Raising awareness about personalization practices is a necessary step toward user-controlled algorithms. Often underestimated in the current discussion is the active capacity of users to make their own choices as to what sites to visit, which people to meet and follow, and what privacy practices to allow.

We conclude that AA goes beyond the technical knowledge of coding and instead involves contextually appreciating and evaluating the FEAT issues behind algorithms. It includes the understanding of how algorithms choose content fairly, how users can engage with algorithms meaningfully, how the awareness of the intent/goals contributes to designing algorithms, and how users are able to reasonably control their data and privacy. Understanding AA will be vital for predicting users' future interests for better performance. The model presented in this study provides insights into ways to incorporate FEAT with awareness and trust features and self-disclosure. Facilitating the adoption of algorithms and enabling trust requires a user perspective of developing understandable/comprehensible AI (Chatterjee et al., 2021), which allows users to meaningfully control AA accordingly.

CRedit authorship contribution statement

Donghee Shin: Principal Investigator, Theoretical Foundations, Discussion, Writing – original draft, Writing – review & editing. **Kerk Kee:** Conceptualization, Theoretical shaping, Data collection process, Data analysis, Literature review on privacy, Discussion, Final editing. **Emily Shin:** Data collection, analysis, literature review, writing, and English editing.

Acknowledgements

This project has been funded by the Research Office of Zayed University, Provost's Research Fellowship Award (R21050/2022) and Research Incentive Fund (R20082/2020).

Appendix. Measurements

Variables	Measures	Sources
Fairness	1. An algorithmic platform does not discriminate against people and does not show favoritism (Nondiscrimination) 2. The source of data throughout an algorithmic process and its data analysis should be accurate and correct (Accuracy) 3. An algorithmic platform complies with the due process requirements of impartiality with no bias (Due process)	(Shin & Park, 2019; Shin, 2021)
Accountability	1. An algorithmic platform requires the person in charge to be accountable for its adverse individual or societal effects in a timely manner (Responsibility) 2. The platforms should be designed to enable third parties to audit and review the behavior of an algorithm (Auditability) 3. The platforms should have the autonomy to change the logic in their entire configuration using only simple manipulations (Controllability)	(Diakopoulos, 2016; Shin & Park, 2019; Shin, 2021)
Transparency	1. The assessment and the criteria of algorithms used should be publicly open and understandable to users (Understandability) 2. Any results generated by an algorithmic system should be interpretable to the users affected by those outputs (Interpretability) 3. Algorithms should let people know how well internal states of algorithms can be understood from knowledge of their external outputs (Observability)	(Diakopoulos, 2016; Rader et al., 2018; Shin, 2021)
Explainability	1. I found algorithmic platforms to be comprehensible. 2. The AI algorithmic services are understandable. 3. I can understand and make sense of the internal workings of personalization.	(Shin, 2021)
Self-disclosure	1. I am planning to share my information with personalized algorithms. 2. I am going to disclose my information to algorithms. 3. I will keep updating my information in algorithmic platforms.	(Sundar et al., 2020; Wang et al., 2016)
Privacy Concerns	1. There is much uncertainty related to giving my personal information to the algorithmic platform. 2. There would be a high risk for loss associated with disclosing my personal information to the platform. 3. The personal information disclosed on the platform is subject to many threats.	(Fast & Jago, 2020)
Trust	1. I trust the recommendations by algorithm-driven platforms.	(Rai, 2020; Renijith et al., 2020;)

(continued on next page)

(continued)

Variables	Measures	Sources
Self-Efficacy	2. Recommended results via algorithmic processes are credible.	(Danielsiek et al., 2017; Hasan et al., 2019; Hu et al., 2021; Mullins & Cronan, 2021)
	3. The algorithmic personalized results are dependable and trustworthy.	
	1. I can perform most of the plans that I have established for myself as to algorithmic platforms.	
	2. When facing difficult situations, I am positive that I can accomplish them through algorithmic platforms.	
	3. I am certain that I can work effectively on different tasks in my interactions with algorithmic platforms.	

References

- Acquisti, A., Brandimarte, L., & Loewenstein, G. (2020). Secrets and likes: The drive for privacy and the difficulty of achieving it in the digital age. *Journal of Consumer Psychology*, 30(4), 736–758.
- Ahmad, F., Widén, G., & Huvila, I. (2020). The impact of workplace information literacy on organizational innovation. *International Journal of Information Management*, 51, Article 102041.
- Akter, S., McCarthy, G., Sajib, S., Michael, K., Dwivedi, Y. K., D'Ambra, J., & Shen, K. (2021). Algorithmic bias in data-driven innovation in the age of AI. *International Journal of Information Management*, 60, Article 102387.
- Alter, S. (2021). Understanding artificial intelligence in the context of usage. *International Journal of Information Management*, Article 102392.
- Ashok, M., Madan, R., Joha, A., & Sivarajah, U. (2022). Ethical framework for artificial intelligence and digital technologies. *International Journal of Information Management*, 62, Article 102433.
- Bandura, A. (1986). *Social foundations of thought and action: A social cognitive theory*. New York: Prentice-Hall.
- Bentler, P. (1990). Comparative fit indexes in structural models. *Psychological Bulletin*, 107(2), 238–246.
- Bollen, K. (1989). A new incremental fit index for general structural equation models. *Sociological Methods & Research*, 17(3), 303–316.
- Borges, A. F., Laurindo, F., Spínola, M., Gonçalves, R., & Mattos, C. (2022). The strategic use of artificial intelligence in the digital era. *International Journal of Information Management*, 57, Article 102225.
- Chatterjee, S., Rana, N., Dwivedi, Y. K., & Baabdullah, A. (2021). Understanding AI adoption in manufacturing and production firms using an integrated TAM-TOE model. *Technological Forecasting and Social Change*, 170, Article 120880.
- Chen, H., & Chen, W. (2015). Couldn't or wouldn't? *Cyberpsychology, Behavior, and Social Networking*, 18(1), 13–19.
- Cotter, K., & Reisdorf, B. (2020). Algorithmic knowledge gaps. *International Journal of Communication*, 14, 745–765.
- Courtois, C., & Timmermans, E. (2018). Cracking the tinder code. *Journal of Computer-Mediated Communication*, 23(1), 1–16.
- Culnan, M., & Armstrong, P. (1999). Information privacy concerns, procedural fairness, and impersonal trust. *Organization Science*, 10, 104–115.
- Danielsiek, H., Toma, L., & Vahrenhold, J. (2017). An instrument to assess self-efficacy in introductory algorithms courses. In *Proceedings of the ACM conference on international computing education research* (pp. 257–65).
- Diakopoulos, N. (2016). Accountability in algorithmic decision making. *Communications of ACM*, 59(2), 58–62.
- Dwivedi, Y. K., Hughes, L., Ismagilova, E., Aarts, G., Coombs, C., Crick, T., ... Williams, M. D. (2021). Artificial Intelligence (AI): Multidisciplinary perspectives on emerging challenges, opportunities, and agenda for research, practice and policy. *International Journal of Information Management*, 57, Article 101994.
- Eslami, M., Rickman, A., Vaccaro, K., Aleyasen, A., Vuong, A., Karahalios, K., & Sandvig, C. (2015). I always assumed that I wasn't really that close to her. In *Proceedings of the 33rd annual ACM conference on human factors in computing systems* (pp. 153–62).
- Fast, N., & Jago, A. S. (2020). Privacy matters or does it? Algorithms, rationalization, and the erosion of concern for privacy. *Current Opinion in Psychology*, 31, 44–48.
- Gran, A., Booth, P., & Bucher, T. (2021). To be or not to be algorithm aware. *Information, Communication & Society*, 24(12), 1779–1796.
- Gruber, J., Hargittai, E., Karaoglu, G., & Brombach, L. (2021). Algorithm awareness as an important internet skill. *International Journal of Communication*, 15, 1770–1788.
- Gutierrez, A., O'Leary, S., Rana, N., Dwivedi, Y. K., & Calle, T. (2019). Using privacy calculus theory to explore entrepreneurial directions in mobile location-based advertising. *Computers in Human Behavior*, 95, 295–306.
- Hair, J., Sarstedt, M., Hopkins, L., & Kuppelwieser, V. (2014). Partial least squares structural equation modeling. *European Business Review*, 26(2), 106–121.
- Hamilton, K., Karahalios, K., Sandvig, C., & Eslami, M. (2014). A path to understanding the effects of algorithm awareness. *CHI '14 Extended Abstracts on Human Factors in Computing Systems*, 631–642.
- Hargittai, E., Gruber, J., Djukaric, T., Fuchs, J., & Brombach, L. (2020). Black box measures? *Information Communication & Society*, 23(5), 764–775.
- Hasan, S., Ahmadi, H., Mortimer, G., Lings, I., Kelly, L., & Kim, H. (2019). Online repurchasing. *Journal of Consumer Affairs*, 54(1), 198–226.
- Hoelter, J. (1983). The analysis of covariance structures. *Sociological Methods & Research*, 11, 325–344.
- Hu, L., & Bentler, P. (1999). Cutoff criteria for fit indexes in covariance structure analysis. *Structural Equation Modeling*, 6(1), 1–55.
- Hu, Q., Lu, Y., Pan, Z., Gong, Y., & Yang, Z. (2021). Can AI artifacts influence human cognition? *International Journal of Information Management*, 56, Article 102250.
- Ischen, C., Araujo, T., Voorveld, H., van Noort, G., & Smit, E. (2020). Privacy concerns in chatbot interactions. In A. Følstad, T. Araujo, S. Papadopoulos, E.L.-C. Law, O.-C. Granmo, E. Luger, & P.B. Brandtzaeg (Eds.), *Chatbot research and design: Third international workshop, Amsterdam, The Netherlands, 2019* (pp. 34–48).
- Jain, S., Basu, S., Dwivedi, Y. K., & Kaur, S. (2022). Interactive voice assistants: Does brand credibility assuage privacy risks? *Journal of Business Research*, 139, 701–717.
- Kizilcec, R. (2016). *How much information?* In *Proceedings of the CHI 2016, 2016*. San Jose, CA.
- Klawitter, E., & Hargittai, E. (2018). It's like learning a whole other language. *International Journal of Communication*, 12, 3490–3510.
- Koenig, A. (2020). The algorithms know me and I know them. *Computers and Composition*, 58, Article 102611.
- Lau, J., Zimmerman, B., & Schaub, F. (2018). Alexa, are you listening? *Proceedings of the ACM on Human-Computer Interaction*, 2, 1–31.
- Lee, T., Lee, B., & Lee-Geiller, S. (2020). The effects of information literacy on trust in government websites. *International Journal of Information Management*, 52, Article 102098.
- Min, S. (2019). From algorithmic disengagement to algorithmic activism. *Telematics and Informatics*, 43, Article 101251.
- Monzer, C., Moeller, J., Helberger, N., & Eskens, S. (2020). User perspectives on the news personalization process. *Digital Journalism*, 8(9), 1142–1162.
- Mullins, J. K., & Cronan, T. (2021). Enterprise systems knowledge, beliefs, and attitude. *International Journal of Information Management*, 59, Article 102348.
- Nishant, R., Kennedy, M., & Corbett, J. (2020). Artificial intelligence for sustainability. *International Journal of Information Management*, 53, Article 102104.
- Preacher, K., & Hayes, A. (2008). Asymptotic and resampling strategies for assessing and comparing indirect effects in multiple mediator models. *Behavior Research Methods*, 40(3), 879–891.
- Rader, E., Cotter, K., & Cho, J. (2018). Explanations as mechanisms for supporting algorithmic transparency. In *Proceedings of the CHI 2018, 2018*. Montréal, QC, Canada.
- Rai, A. (2020). Explainable AI. *Journal of the Academy of Marketing Science*, 48, 137–141.
- Reisdorf, B., & Blank, G. (2020). Algorithmic literacy and platform trust. In E. Hargittai (Ed.), *Handbook of digital inequality*. Edward Elgar Publishing.
- Renjith, S., Sreekumar, A., & Jathavedan, M. (2020). An extensive study on the evolution of context-aware personalized travel recommender systems. *Information Processing & Management*, 57(1), 102078–12.
- Schwartz, S., & Mahnke, M. (2021). Facebook use as a communicative relation. *Information Communication & Society*, 24(7), 1041–1056.
- Shin, D. (2021). The effects of explainability and causability on perception, trust, and acceptance. *International Journal of Human-Computer Studies*, 146, Article 102551.
- Shin, D., & Park, Y. (2019). Role of fairness, accountability, and transparency in algorithmic affordance. *Computers in Human Behavior*, 98, 277–284.
- Shin, D., Zhong, B., & Biocca, F. (2020). Beyond user experience. *International Journal of Information Management*, 52, 102061–102348.
- Siles, I., Segura-Castillo, A., Solís-Quesada, R., & Sancho, M. (2020). Folk theories of algorithmic recommendations on Spotify. *Big Data & Society*, 7(1), 1–15.
- Spanaki, K., Karafili, E., & Despoudi, S. (2021). AI applications of data sharing in agriculture 4.0: A framework for role-based data access control. *International Journal of Information Management*, 59, Article 102350.
- Sundar, S., Kim, J., Beth-Oliver, M., & Molina, M. (2020). Online privacy heuristics that predict information disclosure. In *Proceedings of the CHI'20, 2020*.
- Swart, J. (2021). *Experiencing algorithms*. Social Media + Society. <https://doi.org/10.1177/20563051211008828>.
- Wang, T., Duong, T., & Chen, C. (2016). Intention to disclose personal information via mobile applications. *International Journal of Information Management*, 36(4), 531–542.
- Zarouali, B., Boerman, S., & Vreese, C. (2021). Is this recommended by an algorithm? *Telematics and Informatics*, 62, Article 101607.