



Explaining Loyalty to Generative AI: Extending the Theory of Planned Behavior with AI Anxiety and Perceived Risk

Ibrahim Arpacı^{a,b} , Azize Sahin^c , Ekrem Tatoglu^{d,e}  and Aysun Sahin^f 

^aDepartment of Computer Engineering, Faculty of Engineering, Bursa Uludag University, Bursa, Türkiye; ^bDepartment of Computer Science and Engineering, College of Informatics, Korea University, Seoul, Republic of Korea; ^cDepartment of Business Law, School of Business, Istanbul University, Istanbul, Türkiye; ^dDepartment of Business Administration, College of Business Administration, Gulf University for Science and Technology, Mishref, Kuwait; ^eDepartment of Management, School of Business, Ibn Haldun University, Istanbul, Türkiye; ^fDepartment of Electronics and Automation, Düzce Vocational School, Düzce University, Düzce, Türkiye

ABSTRACT

As generative artificial intelligence (GenAI) use grows, continuance intention does not necessarily imply loyalty. Drawing on the Theory of Planned Behavior (TPB), this study distinguishes continuance intention from loyalty, operationalized as affective attachment and personal significance, and examines AI anxiety and perceived risk as predictors of attitude. Competing deterrence and contextual salience expectations are specified for the risk–attitude relationship. Using PLS-SEM, we analyze data from 1,050 Generation Z users in Turkey in an AI-assisted travel-planning context. AI anxiety relates negatively to attitude, whereas perceived risk relates positively, supporting its predicted direction without establishing contextual salience as the underlying process. Both effects are modest and explain little variance in attitude. Attitude shows the strongest relationship with continuance intention among TPB predictors, while continuance intention is strongly linked to loyalty. The study advances GenAI post-adoption research by distinguishing continuance intention from loyalty and contrasting the roles of AI anxiety and perceived risk.

KEYWORDS

Generative artificial intelligence; user loyalty; AI anxiety; perceived risk; theory of planned behavior

1. Introduction

Generative AI (GenAI) has attracted widespread use across individual and organizational settings, yet adoption alone does not indicate whether users report post-adoption loyalty (Deloitte, 2025; McKinsey & Company, 2025). This distinction has become increasingly important in a market where a growing number of competing GenAI services offer overlapping capabilities and where technical differentiation alone may not sustain platform commitment. In markets characterized by low switching costs and rapid platform proliferation, user reach provides limited evidence of platform-specific loyalty (Farrell & Klemperer, 2007; Rochet & Tirole, 2003). The relevant question is whether repeated use is accompanied by platform-specific loyalty that extends beyond an intention to continue use (Dick & Basu, 1994; Oliver, 1999).

However, information systems (IS) research most directly relevant to GenAI post-adoption has primarily examined acceptance and continuance, with less attention to platform-specific loyalty (Mishra et al., 2023; Vatanasombut et al., 2008). This limitation is particularly relevant in GenAI contexts, where users may continue using several platforms without reporting strong loyalty to any one of them.

Addressing this gap, the present study examines the associations among users' evaluative responses, continuance intention, and platform loyalty. The analysis is situated in the context of AI-assisted travel planning, an information-intensive and consequence-sensitive setting in which users evaluate AI-generated recommendations under uncertainty. Inaccurate or unreliable recommendations in this context may carry financial, temporal, and experiential consequences.

Explaining loyalty in GenAI contexts requires attention not only to functional evaluations but also to psychological frictions associated with AI use. Established frameworks such as the Technology Acceptance Model (TAM) (Davis, 1989) and the Expectation-Confirmation Model (Bhattacharjee, 2001) provide well-developed accounts of acceptance and continuance, but these outcomes alone do not establish loyalty (Li et al., 2026; Venkatesh et al., 2003). GenAI introduces additional evaluative complexity because users may hold favorable and unfavorable evaluations of AI at the same time (Dwivedi et al., 2023; Mick & Fournier, 1998; Schepman & Rodway, 2023). AI anxiety refers to apprehension and unease associated with interacting with AI systems (Wang & Wang, 2022). Recent studies document AI anxiety and AI-related worries in educational and clinical settings (Ayed et al., 2025; Ayed et al., 2025; Batran et al., 2025; Cengiz & Peker, 2025; Zhang & Tong, 2025). Related research on informatics competence, self-efficacy, and decision-making provides complementary grounding for perceived behavioral control in technology-mediated settings (Amer et al., 2025; Batran et al., 2022; Compeau & Higgins, 1995).

The study treats loyalty as a post-adoption outcome distinct from continuance intention and conceptualizes it as a form of affective and preferential commitment (Oliver, 1999). The empirical form of loyalty examined here is narrower, emphasizing affective attachment and personal significance rather than cognitive or behavioral loyalty. It also does not directly capture comparative preference or resistance to alternatives, and the findings should not be generalized to these unmeasured forms of loyalty. This bounded interpretation accords with research distinguishing attitudinal loyalty from retention or repeated behavior (Back & Parks, 2003; Dick & Basu, 1994; Gustafsson et al., 2005; Şahin et al., 2013). Studies of AI-enabled services also examine sustained use, trust, and relational engagement (Boulus-Rødje et al., 2024; Foroughi et al., 2026; Lee et al., 2021, 2023); the present study does not, however, treat these outcomes as interchangeable with loyalty.

The study extends the Theory of Planned Behavior (TPB) (Ajzen, 1991) by examining AI anxiety and perceived risk as distinct psychological frictions associated with attitude. Perceived risk concerns assessments of possible adverse outcomes (Bauer, 1967; Featherman & Pavlou, 2003; Mitchell, 1999), whereas attitude represents the user's overall evaluation of using the technology (Eagly & Chaiken, 1993). Because perceived risk and attitude represent distinct evaluations, a user may recognize risks while still evaluating a platform favorably; such coexistence does not establish a positive association between the constructs. The study therefore compares a deterrence expectation with a contextual salience expectation for the risk-attitude relationship. The contextual salience expectation provides one possible account of why users for whom AI-assisted travel planning is more consequential may recognize platform risks more readily while holding favorable attitudes on the basis of other considerations. Research on AI-chatbot acceptance similarly indicates that evaluative responses vary with users' certainty about their needs (Zhu et al., 2022). However, contextual salience, perceived benefit, usefulness, engagement, and perceived competence are not measured; the comparison evaluates only the direction of the association and does not identify the process underlying it.

The empirical analysis focuses on active Generation Z users of GenAI in Turkey within the context of AI-assisted travel planning. Generation Z provides a relevant population because GenAI is increasingly present in younger adults' digital and working lives, although generational membership is not treated as evidence of uniform technological competence (Deloitte, 2025; Needleman, 2023, 2026; Prensky, 2001). Turkey's ongoing digitalization provides additional contextual grounding (Fidan, 2024).

This study advances research on GenAI post-adoption in three major ways. First, it distinguishes continuance intention from loyalty and clarifies that the empirical construct primarily reflects affective attachment and personal significance within a broader affective and preferential understanding of loyalty (Back & Parks, 2003; Oliver, 1999). Second, it extends TPB by examining AI anxiety and perceived risk as psychological frictions associated with attitude while framing them as a partial rather than comprehensive account of variation in attitude. Finally, it compares deterrence and contextual salience expectations for the perceived risk-attitude relationship, with the comparison limited to the direction of the association rather than constituting a test of contextual salience itself.

2. Theoretical background and hypotheses

Research on digital technology adoption has moved from explaining initial acceptance to understanding users' continued engagement with digital systems after adoption. The Technology Acceptance Model (TAM) (Davis, 1989), the Unified Theory of Acceptance and Use of Technology (UTAUT) (Venkatesh et al., 2003), its consumer-oriented extension, UTAUT2, which incorporates hedonic motivation and habit (Venkatesh et al., 2012), and the Expectation-Confirmation Model of IS continuance (Bhattacharjee, 2001) have each contributed to this trajectory. Collectively, these models show that perceived usefulness, social influence, ease of use, and expectation confirmation are associated with users' decisions to adopt and continue using digital systems. These frameworks have also been widely applied to AI-enabled contexts, including conversational agents (Ashfaq et al., 2020; Blut et al., 2021) and generative systems (Dwivedi et al., 2021, 2023).

This body of research provides a robust theoretical explanation of technology acceptance and continuance. However, these models do not explicitly theorize loyalty as a distinct post-adoption outcome, an omission that becomes especially consequential in the context of generative AI. Continuance intention captures users' stated willingness to keep using a system, whereas loyalty denotes a stronger post-adoption evaluation involving affective attachment and preferential commitment (Bhattacharjee, 2001; Dick & Basu, 1994; Oliver, 1999). Although emerging GenAI research has begun to examine loyalty-related outcomes, the relationship between continuance intention and the specific form of loyalty expressed toward a focal platform remains insufficiently explained.

The distinction between continuance and loyalty is structural rather than terminological. Continuance intention, as Bhattacharjee (2001) operationalizes it, captures an undifferentiated willingness to continue using a system; it does not distinguish users who return out of preference from those who return due to inertia or the absence of viable alternatives. Loyalty, by contrast, reflects a stronger form of commitment grounded in enduring preference and affective attachment (Dick & Basu, 1994; Oliver, 1999). Oliver's (1999) four-stage model makes this progression explicit, proposing that loyalty advances through cognitive, affective, conative, and action phases, each progressively more resistant to disruption. Explaining loyalty therefore requires extending existing post-adoption frameworks beyond continuance intention.

Conceptually, loyalty may encompass affective attachment, preferential commitment, resistance to alternatives, and behavioral persistence (Dick & Basu, 1994; Oliver, 1999). The empirical construct examined in the present study is narrower. As operationalized by the retained indicators, loyalty reflects affective attachment to the focal GenAI platform and the personal significance that users attribute to it. It does not directly capture comparative preference, resistance to alternative platforms, cognitive loyalty, or behavioral loyalty. Accordingly, the construct should be interpreted as a bounded affective-evaluative form of post-adoption loyalty rather than as a comprehensive representation of loyalty.

The importance of this distinction becomes particularly salient in generative AI markets. GenAI tools and platforms are proliferating rapidly, expanding the range of alternatives available to users (Wessel et al., 2025). Where switching costs are relatively low, continuance intention toward a given platform may coexist with parallel or task-specific use of alternatives and may consequently provide limited evidence of loyalty to a focal service (Burnham et al., 2003; Mishra et al., 2023). Users may intend to continue using several platforms without developing a clear affective attachment or preferential orientation toward any one of them. Loyalty therefore warrants examination as a distinct post-adoption construct rather than being inferred from continuance intention alone (Dick & Basu, 1994; Oliver, 1999).

2.1. Theory of planned behavior

Among the available frameworks, the TPB (Ajzen, 1991) provides a particularly suitable foundation for explaining post-adoption evaluations in the GenAI context. Within this framework, behavioral intention is shaped by three proximal determinants: subjective norms, attitude, and perceived behavioral control (PBC). Attitude captures the overall evaluative response (Eagly & Chaiken, 1993). Subjective norms capture perceived social pressure from relevant others, and PBC reflects perceived feasibility given available abilities and constraints, with self-efficacy as a central component (Compeau & Higgins, 1995;

Eagly & Chaiken, 1993). TPB has been applied in contexts in which evaluation, social influence, and perceived competence are jointly relevant to behavioral decisions (Glasman & Albarracín, 2006; Taylor & Todd, 1995; Venkatesh et al., 2003). The framework has also been extended to incorporate additional evaluative and risk-related constructs in electronic commerce and digital technology settings closely related to the present study (Pavlou et al., 2007).

TPB is well-suited to examining GenAI use because its core constructs capture evaluation, social influence, and perceived behavioral control in contexts where users must assess both technological value and personal capability. In the present context, attitude refers to the user's overall evaluation of using a GenAI platform, including whether continued interaction with it is regarded favorably. In GenAI settings, this evaluation may be complicated by outputs that users cannot always verify readily. Subjective norms capture the increasing social legitimacy of AI use in academic and professional environments (Dwivedi et al., 2021; Ivanov et al., 2024). Perceived behavioral control reflects whether users feel capable of prompting effectively, interpreting outputs critically, and managing system opacity (Compeau & Higgins, 1995; Long & Magerko, 2020). However, the original framework does not fully capture the psychological frictions that distinguish GenAI from conventional digital tools.

Research on responses to algorithmic systems further shows that evaluations of these systems vary with observed performance, task characteristics, and user conditions. Dietvorst et al. (2015) found that observing an algorithm make errors can reduce subsequent reliance even when its forecasts remain more accurate than human forecasts. Related research identifies task characteristics and trust-relevant conditions as sources of variation in responses to algorithmic systems (Castelo et al., 2019; Glikson & Woolley, 2020). Under other conditions, Logg et al. (2019) document algorithm appreciation, in which users assign greater weight to algorithmic than to human advice, while also identifying variation across users and comparison conditions. These findings do not establish a perceived-risk explanation, but they show that responses to algorithmic systems are sensitive to the conditions under which users encounter and evaluate them.

This context sensitivity indicates that TPB's original specification may not fully represent the psychological concerns associated with GenAI use. Extensions of TPB are warranted when domain-specific constructs account for aspects of intention that are not adequately represented by its original determinants (Ajzen, 2020; Conner & Armitage, 1998). Accordingly, the present study incorporates AI anxiety and perceived risk as distinct psychological frictions associated with attitude and positions loyalty as a post-adoption outcome associated with continuance intention.

2.2. AI anxiety

Within TPB, attitude represents the evaluative component most proximally associated with behavioral intention (Ajzen, 1991, 2001). In the present model, AI anxiety and perceived risk are treated as psychologically distinct antecedents of attitude. AI anxiety represents an affective response characterized by apprehension and discomfort when engaging with or anticipating AI systems (Schepman & Rodway, 2023; Tong et al., 2025; Wang & Wang, 2022). Perceived risk represents a cognitive appraisal of the probability and magnitude of adverse outcomes associated with use (Bauer, 1967; Featherman & Pavlou, 2003; Lim, 2003). Although both reflect broader orientations toward AI systems, their evaluative consequences are assessed at the focal-platform level, consistent with the conceptualization of attitude, continuance intention, and loyalty in this study. Because these frictions are psychologically distinct, their associations with attitude should be theorized and tested separately.

AI anxiety is particularly relevant in GenAI contexts because system opacity and output uncertainty can sustain affective discomfort during interaction. When interaction evokes unease, users' evaluative attention may shift from perceived benefits such as efficiency, creativity, and task support to concerns about overreliance, reduced autonomous judgment, and inaccurate information. Within TPB, this affective reorientation can weaken the evaluative beliefs that underlie the formation of favorable attitudes (Ajzen, 2001). Consequently, users experiencing greater AI anxiety are expected to evaluate GenAI systems less favorably. This weaker evaluative foundation is, in turn, expected to be associated with lower continuance intention and ultimately weaker loyalty.

Prior research in AI-enabled work, learning, and service contexts supports this expectation (Chen et al., 2025; Li et al., 2026; Wang et al., 2022). Evidence from emerging GenAI research provides more direct support for this expectation. Cengiz and Peker (2025) show that AI anxiety is negatively related to acceptance of generative AI through attitudes toward AI and AI literacy, while Zhang and Tong (2025) link ChatGPT-related technostress to AI anxiety, negative attitudes, and discontinuance-oriented responses. Evidence from clinical and professional AI-use contexts further indicates that AI-related worries, informatics competence, and self-efficacy become salient when users rely on AI systems in consequence-sensitive decision environments (Amer et al., 2025; Ayed et al., 2025; Ayed et al., 2025; Batran et al., 2022, 2025). These findings are especially relevant to AI-assisted travel planning, where users must evaluate probabilistic recommendations and make decisions that may have financial, temporal, and experiential consequences. This reasoning and prior evidence support the expectation that AI anxiety is negatively associated with attitude toward generative AI. Accordingly:

H1. AI anxiety is negatively associated with attitude toward generative AI.

2.3. Competing directional expectations for perceived risk

Perceived risk refers to users' assessments of uncertainty and possible adverse consequences associated with technology use (Bauer, 1967; Featherman & Pavlou, 2003). In the present study, the construct encompasses concerns about personal-data use, harm to privacy or confidentiality, general uncertainty, loss of control over data, and harmful or misleading outputs, reflecting privacy- and output-related risks documented in AI contexts (Gerlich, 2024; Weidinger et al., 2022). Related evidence from consequence-sensitive healthcare settings also documents AI-related worries involving data management and regulatory or ethical concerns, although these broader worries are not equivalent to the perceived-risk construct examined here (Ayed et al., 2025; Batran et al., 2025). Perceived risk is generally treated as a deterrent in digital-service research, where previous studies report negative associations with adoption-related evaluations and intentions (Featherman & Pavlou, 2003; Kim et al., 2008). This evidence provides a basis for expecting users who perceive greater risk to hold less favorable attitudes toward GenAI.

Perceived risk and attitude are nevertheless distinct evaluations. Perceived risk concerns possible adverse outcomes, whereas attitude represents the user's overall evaluation of using the technology (Eagly & Chaiken, 1993). Privacy-calculus research examines privacy concerns alongside other considerations associated with willingness to transact online (Dinev & Hart, 2006), while consumer-risk research links perceived risk to information search and more extensive evaluation (Mitchell, 1999). Neither line of research shows that perceived risk improves attitude or is positively associated with it. These literatures do, however, indicate that risk recognition can remain relevant alongside other evaluations. A user may therefore recognize privacy-, data-control-, or output-related risks while still holding a favorable overall evaluation of a GenAI platform. This coexistence does not establish a positive association, which requires perceived risk and attitude to covary systematically across users.

One possible theoretical basis for such covariance concerns differences in the subjective salience of the use context. Although all respondents evaluated GenAI in the context of travel planning, the financial, temporal, or experiential consequences of relying on AI-generated recommendations may be more salient to some users than to others. Users for whom this context is more consequential may recognize platform risks more readily while also evaluating the platform favorably on the basis of other, unmeasured considerations. Such heterogeneity could be reflected in positive covariance between perceived risk and attitude without implying that greater perceived risk itself produces a more favorable attitude. Research on algorithm aversion and appreciation does not examine perceived risk directly, but it shows that evaluations of algorithmic systems can vary across tasks, users, and decision conditions (Dietvorst et al., 2015; Logg et al., 2019). We refer to this account as contextual salience.

This account cannot be evaluated directly, however, because the model contains no measures of perceived benefit, usefulness, engagement, or perceived competence. Limited risk awareness does not straightforwardly explain a positive association under the present operationalization, because lower awareness of the measured risks would ordinarily be reflected in lower perceived-risk scores among

favorably disposed users. Such a pattern would tend to weaken rather than generate positive covariance. However, because objective risk knowledge was not assessed, this alternative cannot be conclusively evaluated. The screening criteria establish active GenAI use, but they neither demonstrate risk literacy nor exclude early-adopter characteristics. Perceived competence also remains a directionally indeterminate omitted variable because it could accompany both stronger risk recognition and more favorable attitudes or, alternatively, reduce perceived vulnerability while remaining positively associated with attitude. Other unmeasured characteristics cannot be ruled out. Because contextual salience is not measured, a positive coefficient would support only the direction of the association and would neither show that perceived risk makes attitudes more favorable nor constitute a test of this account. The contextual salience label therefore identifies the theoretical account informing the positive directional expectation rather than a process tested by the model. We therefore compare two directional expectations without treating either explanation as established:

H2a. Perceived risk is negatively associated with attitude toward generative AI (Deterrence expectation).

H2b. Perceived risk is positively associated with attitude toward generative AI (Contextual salience expectation).

2.4. Attitude

Within the TPB framework, attitude is positioned as a proximal antecedent of behavioral intention, such that favorable evaluation is expected to be associated with stronger intention to continue the behavior (Ajzen, 1991; Bhattacharjee, 2001; Davis, 1989). In the GenAI context, initial use may reflect curiosity, experimentation, or external pressure rather than a settled evaluation of the focal platform. Regardless of what prompts initial use, users who evaluate continued interaction more favorably are expected to report stronger continuance intention.

The role of attitude is particularly relevant in decision contexts where reliance on AI outputs carries meaningful consequences. Users must assess whether AI-generated recommendations are sufficiently useful and reliable to support their choices. At this stage, attitude serves as a central evaluative basis for continuance intention rather than a direct causal determinant of continued use. Prior evidence from ChatGPT and generative AI contexts is consistent with the positive association between favorable attitudes and continuance intention (Bhattacharjee, 2001; Han et al., 2024; Pavlou & Fygenson, 2006). Accordingly:

H3. Attitude is positively associated with continuance intention toward generative AI.

2.5. Subjective norms

Within TPB, behavioral intention is associated not only with individuals' own evaluations but also with their perceptions of how relevant others view the behavior. Subjective norms refer to perceived social pressure from relevant others and may be especially influential when expectations concerning appropriate behavior remain unsettled (Ajzen, 1991; Venkatesh & Davis, 2000). In the GenAI context, norms governing appropriate use are still evolving across academic and professional environments, making users more likely to rely on social cues when judging whether continued use is appropriate.

This reasoning is consistent with social comparison processes and the role of interpersonal influence in technology diffusion (Festinger, 1954; Rogers, 2003). Endorsement from peers, instructors, colleagues, and supervisors may signal that continued GenAI use is socially acceptable. Subjective norms are therefore expected to be positively associated with users' willingness to continue using the focal platform, although the present model does not measure perceived legitimacy or social uncertainty as separate explanatory processes.

The role of subjective norms is particularly relevant in contexts where the consequences of reliance on AI outputs are non-trivial. Users may look to peers and other referents to assess whether relying on

GenAI for consequential decisions is appropriate and acceptable. Prior evidence from AI-supported learning and ChatGPT-related contexts is consistent with the positive association between subjective norms and continuance intention (Saxena & Doleck, 2023; Wang et al., 2025). Accordingly:

H4. Subjective norms are positively associated with continuance intention toward generative AI.

2.6. Perceived behavioral control

Perceived behavioral control captures the role of perceived feasibility in the formation of behavioral intention. It reflects users' assessment of whether they possess the resources, competencies, and situational conditions required to perform the behavior successfully, with self-efficacy as its central cognitive component (Ajzen, 1991; Bandura, 1986).

In the GenAI context, perceived behavioral control is particularly relevant because effective use may require iterative interaction, critical evaluation of outputs, and management of system limitations. Users who doubt their ability to perform these activities may report weaker continuance intention even when they evaluate the platform favorably. Conversely, users who regard continued use as feasible and manageable are expected to report stronger continuance intention.

Prior information-systems research reports positive associations of perceived behavioral control and self-efficacy with behavioral intention (Compeau & Higgins, 1995; Venkatesh et al., 2003), and emerging evidence from GenAI contexts indicates a similar pattern (Mun & Hwang, 2025). Research on informatics competency, self-efficacy, and clinical decision-making further supports the relevance of capability-related perceptions in technology-mediated decision settings, although these studies do not test continuance intention toward GenAI platforms directly (Amer et al., 2025; Batran et al., 2022). Accordingly:

H5. Perceived behavioral control is positively associated with continuance intention toward generative AI.

2.7. Continuance intention and loyalty

The final path in the proposed model concerns the association between continuance intention and loyalty, the focal post-adoption outcome. Oliver's (1999) framework distinguishes cognitive, affective, conative, and action manifestations of loyalty, indicating that intention, attachment, and behavior are not interchangeable. In the present study, continuance intention captures users' stated willingness to continue using the focal platform, whereas the retained loyalty indicators capture affective attachment and the platform's personal significance to the user. The retained indicators do not directly assess comparative preference, resistance to alternatives, cognitive loyalty, or behavioral loyalty. Continuance intention and the measured form of loyalty are therefore treated as related but conceptually distinct post-adoption evaluations rather than as successive stages observed over time.

Research on loyalty and habitual use similarly indicates that repeated behavior and evaluative loyalty are related but distinct outcomes (Dick & Basu, 1994; Limayem et al., 2007). Studies in online and service settings also distinguish continued patronage intentions from attitudinal and behavioral loyalty, supporting the treatment of these outcomes as connected but nonequivalent (Cossío-Silva et al., 2016; Srinivasan et al., 2002). More directly, evidence from mobile payment applications reports a positive association between continuance intention and e-loyalty (Al Amin et al., 2024). Although these studies concern digital services rather than GenAI-assisted travel planning, they provide a basis for expecting continuance intention to be positively associated with the retained form of loyalty examined here. Accordingly:

H6. Continuance intention is positively associated with loyalty toward generative AI.

The proposed model treats AI anxiety and perceived risk as distinct psychological frictions associated with attitude toward the focal GenAI platform. The perceived risk-attitude relationship is specified through two competing directional expectations: a deterrence expectation and a contextual salience

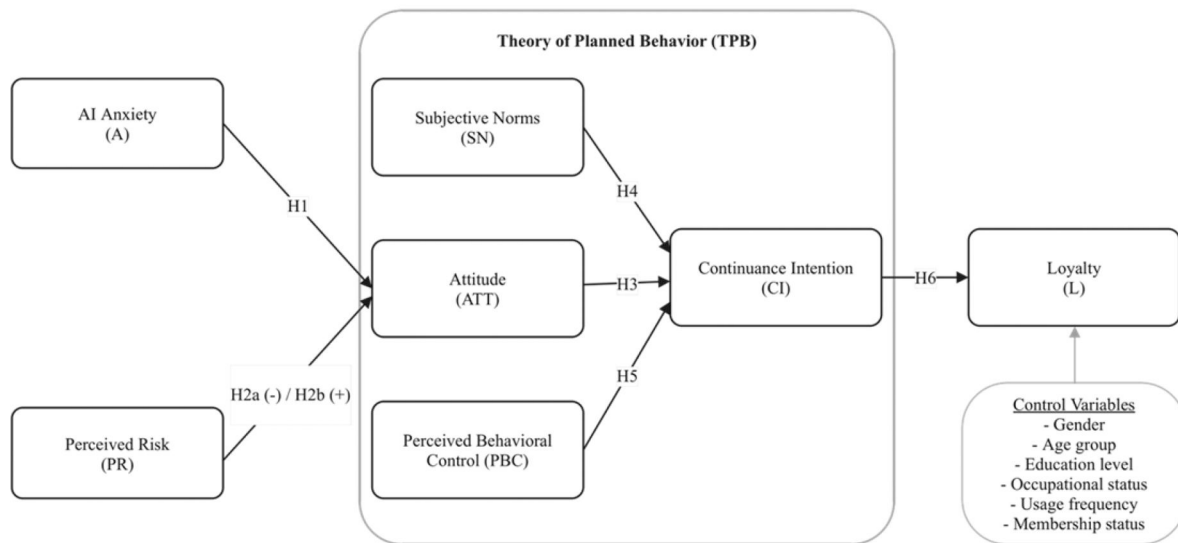


Figure 1. Conceptual model.

Note: H2a and H2b represent competing hypotheses for the perceived risk → attitude path. H2a reflects the deterrence perspective, whereas H2b reflects the contextual salience expectation.

expectation. Attitude, subjective norms, and perceived behavioral control are specified as predictors of continuance intention, which is in turn associated with the retained form of loyalty, operationalized in this study as affective attachment and personal significance. For the perceived risk–attitude relationship, the comparison evaluates which directional expectation is supported; it does not test contextual salience as the process underlying the observed association. The resulting conceptual framework is presented in Figure 1.

3. Research methodology

3.1. Sample and data collection

The study was conducted in Turkey among Generation Z users with recent and substantive experience using generative AI, with the empirical context framed around AI-assisted travel planning. This empirical setting was selected because travel planning involves information-intensive, consequence-sensitive decisions that require evaluating multiple alternatives and assessing the reliability of AI-generated recommendations (Kim et al., 2020; Lalicic & Weismayer, 2021; Li & Lee, 2025). Turkey provides a relevant context as a rapidly digitalizing emerging economy in which AI-related practices are being adopted while social norms, regulatory expectations, and trust arrangements are still developing (Deloitte, 2024). These characteristics make Turkey a relevant setting for examining psychological frictions associated with GenAI evaluation.

Generation Z was selected as the focal population because digital platforms are deeply embedded in their academic, professional, and social routines, and because travel planning is a familiar decision domain for this cohort (Deloitte, 2025; Needleman, 2023). At the same time, familiarity with digital technologies does not eliminate uncertainty, perceived risk, or anxiety in AI-enabled interactions, making this group suitable for examining post-adoption evaluations among frequent users (Popşa, 2024).

Data were collected through an online survey between September and October 2025. Because no comprehensive sampling frame exists for users of GenAI in travel-planning contexts, probability-based sampling was not feasible. A quota-based purposive sampling strategy was therefore used to ensure variation across gender, geographic region, and occupational status, with quotas monitored during data collection to maintain balance across these categories, consistent with post-adoption research in contexts where verified user experience is particularly important (Calder et al., 1981; Palinkas et al., 2015). Accordingly, the sample should be interpreted as a theoretically relevant non-probability sample rather than as a statistically representative sample of all Generation Z GenAI users in Turkey. Compared with

Table 1. Demographic profile of respondents (N = 1,050).

Characteristics	Number	%	Characteristics	Number	%
Gender			Membership status		
Female	584	55.6	Paid subscription	245	23.3
Male	466	44.4	Free access	805	76.7
Age groups			GenAI use for travel planning		
18–23 years old	540	51.4	Every day	362	34.5
24–28 years old	510	48.6	Several times a week	452	43.0
			Several times a month	92	8.8
			Rarely	144	13.7
Education			Gen-AI tools		
High school or below	150	14.3	ChatGPT	999	95.1
Bachelor's degree	700	66.7	Google Gemini	308	29.3
Graduate degree	200	19.0	Microsoft Copilot	159	15.1
Occupation			Claude	56	5.3
Student	457	43.5	Midjourney/DALL-E	47	4.5
Employed	584	55.6	Other	78	7.4
Other	9	0.9			

Note: GenAI tools were measured using a multiple-response item; therefore, frequencies and percentages exceed N = 1,050 and 100%, respectively.

unrestricted convenience recruitment, quota monitoring provided greater control over sample composition by maintaining prespecified variation across gender, geographic region, and occupational status and reducing the likelihood that recruitment would be disproportionately concentrated within particular categories (Mercer et al., 2017). This approach therefore promoted greater heterogeneity across the monitored characteristics without eliminating self-selection or conferring the representativeness associated with probability sampling. Participants were recruited through AI-focused online communities, peer referral networks, and targeted outreach on Instagram and LinkedIn, yielding a sample that extends beyond university-affiliated respondents and includes employed young adults from diverse occupational backgrounds.

To capture a meaningful post-adoption experience, respondents were required to meet four screening criteria based on their general GenAI use: use within the previous 7 days, at least twice-weekly usage, experience with at least two platforms, and engagement in substantive tasks. The questionnaire was framed around AI-assisted travel planning, with key items asking respondents to evaluate their attitudes, intentions, and use of AI tools in that context, so that responses reflected a specific, behaviorally relevant setting rather than general perceptions of AI. Accordingly, the usage-frequency variable reported in Table 1 reflects how often respondents used GenAI for travel planning specifically, whereas eligibility for the study was determined by general GenAI use.

Participation in the study was voluntary and anonymous. Before completing the questionnaire, respondents were informed about the academic purpose of the study, and electronic informed consent was obtained. All responses were used solely for research purposes. The study was conducted in accordance with the research ethics guidelines of Ibn Haldun University and relevant national standards. According to the written institutional policy of Ibn Haldun University, formal ethics committee approval was not required because the study involved anonymous, non-sensitive survey data, collected no personally identifiable information, and included adult respondents who participated voluntarily.

Of the 1,389 initial responses, 339 were excluded based on pre-specified data-quality criteria, including straight-lining, failed attention checks, insufficient platform familiarity, or implausibly short completion times (Ward & Meade, 2023), yielding a final analytic sample of 1,050. This sample size meets commonly cited PLS-SEM recommendations and is adequate for estimating the proposed model (Hair et al., 2019; Hair & Alamer, 2022; Memon et al., 2020; Sarstedt et al., 2022). Sample characteristics are reported in Table 1.

Nonresponse bias was assessed using an early-late respondent comparison (Armstrong & Overton, 1977). Chi-square tests for categorical demographic variables and independent-samples t-tests for focal constructs indicated no significant differences between early and late respondents across the variables assessed (all $p > 0.05$), suggesting that there was no evidence of substantial nonresponse bias based on this diagnostic.

Table 2. Measures.

Constructs	Items
AI anxiety (Appel et al., 2020; Wang & Wang, 2022)	A1: "Using this AI platform makes me feel uneasy." A2: "I feel anxious when I rely on this AI platform." A3: "I feel uncomfortable because I cannot fully understand how this AI platform produces its outputs." A4: "Using this AI platform sometimes makes me feel a loss of control." A5: "I feel uneasy about making decisions based on this AI platform's responses."
Perceived risk (Featherman & Pavlou, 2003)	PR1: "I am concerned about how this AI platform uses my personal data." PR2: "Using this AI platform could potentially harm my privacy or confidentiality." PR3: "I feel that there are uncertainties or potential risks associated with using this AI platform." PR4: "I am worried that I may lose control over my data when interacting with this AI platform." PR5: "There is a possibility that this AI platform may produce harmful or misleading outputs."
Subjective norms (Al-Emran et al., 2020)	SN1: "People who are important to me think I should use this AI platform." SN2: "People who influence my behavior think I should use this AI platform." SN3: "People whose opinions I value prefer that I use this AI platform."
Attitude (Davis, 1989)	ATT1: "Using this AI platform is a good idea." ATT2: "I have a generally favorable attitude toward using this AI platform." ATT3: "Using this AI platform is beneficial for my tasks."
Perceived behavioral control (Ajzen, 2002)	PBC1: "I have sufficient knowledge to use this AI platform." PBC2: "I feel that using this AI platform is under my control." PBC3: "I feel confident in my ability to use this AI platform effectively."
Continuance intention (Al-Sharafi et al., 2023; Venkatesh et al., 2012)	C11: "I am likely to continue using this AI platform in the future." C12: "I intend to continue using this AI platform regularly." C13: "I have a strong intention to keep using this AI platform over time."
Loyalty (Oliver, 1999)	L1: "This AI platform is my first choice among available tools." L2: "I feel emotionally attached to this AI platform." L3: "This AI platform means a lot to me." L4: "Using this AI platform matters to me." L5: "I would continue using this AI platform even if a comparable alternative became available."

Note: A5 was initially measured but was removed because of its very low outer loading. L1 and L5 were removed from the loyalty construct during construct purification to reduce conceptual overlap with continuance intention and switching-resistance-oriented wording.

3.2. Measures

All focal constructs were operationalized using multi-item reflective scales adapted from established measures in the information systems and consumer behavior literature. Item wording was revised to reflect the use of generative AI in travel-planning contexts while preserving the original conceptual meaning of each scale. All constructs were measured using a five-point Likert-type scale ranging from 1 ("strongly disagree") to 5 ("strongly agree").

The questionnaire was originally developed in English and subsequently translated into Turkish using a double-blind back-translation procedure to ensure conceptual equivalence across the two language versions (Brislin, 1986; Douglas & Craig, 2007). Throughout the survey, respondents evaluated the items with reference to their use of generative AI for travel-planning purposes.

AI anxiety (A) was initially measured using five items adapted from Wang and Wang (2022) and Appel et al. (2020), capturing affective responses such as apprehension, discomfort, and unease associated with interacting with AI systems.

Perceived risk (PR) was measured using five items adapted from Featherman and Pavlou (2003). The items assess users' perceptions of potential negative outcomes, with an emphasis on privacy and data-control concerns, together with general uncertainty and the possibility of harmful or misleading outputs.

Subjective norms (SN) were measured using three items adapted from Al-Emran et al. (2020), consistent with TPB-based information systems research (Ajzen, 1991; Taylor & Todd, 1995). The scale captures perceived social pressure from relevant others regarding the use of generative AI.

Attitude (ATT) was measured using three items adapted from Davis (1989). The items reflect users' overall evaluative orientation toward the use of generative AI.

Perceived behavioral control (PBC) was measured using three items adapted from Ajzen (2002). The scale assesses perceived capability to use generative AI effectively, including confidence in managing tasks and interpreting outputs.

Continuance intention (CI) was measured using three items adapted from Venkatesh et al. (2012) and Al-Sharafi et al. (2023). The items capture users' intention to continue using generative AI tools.

Loyalty (L) was initially measured using five items designed with reference to Oliver's (1999) framework and prior research on digital services (Cossío-Silva et al., 2016; Srinivasan et al., 2002). The initial item set covered comparative choice, affective attachment, personal significance, and stated resistance to a comparable alternative. Following the removal of L1 and L5 during construct purification, the retained items (L2–L4) operationalize affective attachment and personal significance toward the focal GenAI platform. They do not directly measure comparative preference, resistance to alternatives, cognitive loyalty, or behavioral loyalty. The findings concerning loyalty should therefore be interpreted within the narrower scope represented by these retained items.

Several control variables were included to account for individual-level variation in generative AI use and post-adoption evaluation. Demographic controls included gender, age group, education level, and occupational status, as these characteristics may be associated with technology adoption and post-adoption responses. Two usage-related controls were also incorporated. Usage frequency was measured categorically to capture variation in the intensity of interaction with generative AI tools. Membership status, distinguishing paid subscription from free access, was included to account for possible differences in access mode and financial commitment associated with post-adoption evaluation. Including these controls allowed the focal structural associations to be estimated while accounting for these individual differences. The full list of measurement items and their corresponding sources is reported in Table 2.

3.3. Common method bias

Common method bias (CMB) was assessed using both procedural and statistical safeguards, given the cross-sectional and self-reported nature of the data (Podsakoff et al., 2012, 2024). Procedurally, respondents were assured of anonymity, informed that there were no right or wrong answers, and presented with predictor and outcome measures in separate sections of the questionnaire to reduce evaluation apprehension and response-pattern bias.

Statistically, Harman's single-factor test showed that the first unrotated factor explained 44.04% of the total variance, below the commonly used 50% benchmark (Harman, 1976). Because Harman's test is a limited diagnostic, full collinearity was also assessed following Kock's (2015) approach. The resulting full-collinearity VIF values ranged from 1.022 to 2.688 and remained below the recommended threshold of 3.3, providing no indication of severe collinearity-based common method concerns under this diagnostic (Kock, 2015; Kock & Lynn, 2012).

Although these procedures reduce concerns regarding CMB, they cannot completely eliminate the possibility of same-source effects. Because all constructs were measured using a single survey instrument administered at one point in time, some degree of CMB may still be present and should therefore be considered when interpreting the findings. Consequently, the relatively strong associations observed among attitude, continuance intention, and loyalty should be interpreted with caution. Because no theoretically unrelated marker variable was incorporated during the questionnaire design stage, and because the variance-based estimation approach used in the main analysis does not readily accommodate an unmeasured latent method factor, marker-variable and unmeasured latent method factor approaches could not be implemented. Future research could strengthen the assessment of CMB by incorporating an appropriate marker variable during questionnaire design.

4. Data analysis and results

The proposed model was estimated using PLS-SEM in SmartPLS 4.1.1.5 (Ringle et al., 2024). PLS-SEM is appropriate for models that combine established theoretical constructs with context-specific antecedents and for research that emphasizes both predictive performance and theory-oriented model assessment (Hair & Alamer, 2022; Sarstedt et al., 2017).

Table 3. Measurement model.

Constructs/items	FL	α	ρ_A	ρ_C	AVE	VIF
AI anxiety (A)	—	0.844	0.893	0.872	0.632	—
A1	0.927	—	—	—	—	1.660
A2	0.771	—	—	—	—	2.039
A3	0.726	—	—	—	—	2.390
A4	0.740	—	—	—	—	1.949
Attitude (ATT)	—	0.853	0.856	0.911	0.773	—
ATT1	0.863	—	—	—	—	2.001
ATT2	0.899	—	—	—	—	2.294
ATT3	0.875	—	—	—	—	2.074
Continuance intention (CI)	—	0.888	0.890	0.931	0.817	—
CI1	0.894	—	—	—	—	2.525
CI2	0.921	—	—	—	—	2.969
CI3	0.898	—	—	—	—	2.405
Loyalty (L)	—	0.841	0.848	0.903	0.757	—
L2	0.859	—	—	—	—	1.742
L3	0.869	—	—	—	—	2.203
L4	0.883	—	—	—	—	2.225
Perceived behavioral control (PBC)	—	0.864	0.872	0.917	0.786	—
PBC1	0.859	—	—	—	—	2.090
PBC2	0.908	—	—	—	—	2.663
PBC3	0.891	—	—	—	—	2.164
Perceived risk (PR)	—	0.883	0.906	0.912	0.677	—
PR1	0.843	—	—	—	—	2.105
PR2	0.865	—	—	—	—	2.628
PR3	0.881	—	—	—	—	2.675
PR4	0.771	—	—	—	—	2.158
PR5	0.745	—	—	—	—	1.680
Subjective norms (SN)	—	0.870	0.874	0.920	0.793	—
SN1	0.892	—	—	—	—	2.253
SN2	0.893	—	—	—	—	2.521
SN3	0.886	—	—	—	—	2.180

Note: FL = factor loading; α = Cronbach's alpha; ρ_A = rho_A; ρ_C = rho_C; AVE = average variance extracted; VIF = variance inflation factor. Construct-level reliability and convergent validity statistics are reported on construct rows, whereas factor loadings and indicator VIF values are reported on indicator rows. Dashes indicate that the statistic is not applicable at that level of reporting. The table reports the retained indicators used in the final measurement model.

Statistical inference was based on nonparametric bootstrapping with 10,000 resamples, using bias-corrected and accelerated confidence intervals. The analysis proceeded in two stages: evaluation of the measurement model, followed by estimation of the structural model. Predictive performance was subsequently assessed using cross-validated predictive ability testing (CVPAT) and PLSpredict (Sharma et al., 2023; Shmueli et al., 2019).

4.1. Measurement model

The reflective measurement model was evaluated following the procedure recommended by Hair et al. (2019), in terms of internal consistency, indicator reliability, convergent validity, and discriminant validity.

Table 3 reports indicator reliability, internal consistency, and convergent validity for the retained measurement model. A5 was removed from the AI anxiety construct because its outer loading fell below the recommended threshold. L1 and L5 were removed from the loyalty construct during construct purification to improve discriminant validity, with L5 also displaying conceptual proximity to continuance intention. Their removal restricted the retained loyalty construct to affective attachment and personal significance, as represented by L2–L4. All retained indicator loadings exceeded 0.70, ranging from 0.726 to 0.927, indicating adequate item reliability. Internal consistency was supported across all constructs. Cronbach's alpha values ranged from 0.841 to 0.888, rho_A (ρ_A) values ranged from 0.848 to 0.906, and composite reliability (ρ_C) values ranged from 0.872 to 0.931, all exceeding the recommended threshold of 0.70 (Hair & Alamer, 2022).

Convergent validity was supported, as average variance extracted (AVE) values ranged from 0.632 to 0.817, each above the recommended minimum of 0.50 (Fornell & Larcker, 1981). Indicator-level variance inflation factor (VIF) values ranged from 1.660 to 2.969, indicating no problematic collinearity.

Discriminant validity was evaluated using the Fornell–Larcker criterion and the heterotrait–monotrait ratio (HTMT). As reported in Table 4, the Fornell–Larcker criterion was satisfied because the

Table 4. Discriminant validity (Fornell–Larcker criterion).

Construct	1	2	3	4	5	6	7
1. AI anxiety (A)	0.795	—	—	—	—	—	—
2. Attitude (ATT)	−0.156	0.879	—	—	—	—	—
3. Continuance intention (CI)	−0.167	0.756	0.904	—	—	—	—
4. Loyalty (L)	−0.114	0.696	0.734	0.870	—	—	—
5. Perceived behavioral control (PBC)	−0.099	0.582	0.603	0.512	0.886	—	—
6. Perceived risk (PR)	0.476	0.115	0.132	0.031	0.155	0.823	—
7. Subjective norms (SN)	0.005	0.484	0.475	0.547	0.330	0.040	0.890

Note: Diagonal elements represent the square root of the AVE for each construct.

Table 5. HTMT ratios and bootstrap confidence intervals.

Construct pair	HTMT	95% CI
ATT – A	0.135	[0.088, 0.189]
CI – A	0.140	[0.083, 0.208]
CI – ATT	0.867	[0.825, 0.903]
L – A	0.100	[0.060, 0.158]
L – ATT	0.812	[0.766, 0.850]
L – CI	0.840	[0.805, 0.870]
PBC – A	0.098	[0.043, 0.168]
PBC – ATT	0.675	[0.608, 0.733]
PBC – CI	0.686	[0.624, 0.738]
PBC – L	0.592	[0.527, 0.650]
PR – A	0.555	[0.497, 0.608]
PR – ATT	0.123	[0.059, 0.206]
PR – CI	0.136	[0.072, 0.216]
PR – L	0.065	[0.039, 0.080]
PR – PBC	0.166	[0.090, 0.246]
SN – A	0.044	[0.018, 0.061]
SN – ATT	0.559	[0.501, 0.612]
SN – CI	0.536	[0.475, 0.592]
SN – L	0.640	[0.581, 0.695]
SN – PBC	0.379	[0.309, 0.447]
SN – PR	0.048	[0.021, 0.102]

Note: A = AI anxiety; ATT = attitude; CI = continuance intention; L = loyalty; PBC = perceived behavioral control; PR = perceived risk; SN = subjective norms. HTMT = heterotrait–monotrait ratio. Confidence intervals are based on bootstrapping. Discriminant validity is supported when the HTMT confidence interval does not include 1.00.

square root of each construct's AVE exceeded its correlations with all other constructs (Fornell & Larcker, 1981). This evidence also supported the distinction between continuance intention and loyalty: the square roots of AVE for continuance intention (0.904) and loyalty (0.870) were both higher than their inter-construct correlation (0.734).

Table 5 presents the HTMT results and bootstrap confidence intervals. Most HTMT values fell below the conservative 0.85 threshold, and every 95% confidence interval excluded 1.00 (Henseler et al., 2015). Two post-adoption construct pairs warranted explicit attention. CI–ATT produced the highest HTMT value (HTMT = 0.867, 95% CI [0.825, 0.903]), exceeding the conservative 0.85 cutoff. L–CI produced a value close to this threshold (HTMT = 0.840, 95% CI [0.805, 0.870]). Because attitude, continuance intention, and loyalty are conceptually related post-adoption constructs, the 0.90 criterion recommended for conceptually related constructs provides a relevant benchmark for these pairs (Henseler et al., 2015). Both values remained below this threshold, and neither confidence interval included 1.00. As noted above, L1 and L5 were removed during scale purification in response to the conceptual and empirical overlap between loyalty and continuance intention, with L5's wording (“I would continue using this AI platform even if a comparable alternative became available”) closely resembling continuance intention. The retained loyalty construct (L2–L4) therefore sharpens this distinction. Thus, discriminant validity was supported while the empirical closeness among these post-adoption constructs was explicitly acknowledged.

4.2. Structural model and hypothesis testing

As reported in Table 6, the model explained 7.1% of the variance in attitude ($R^2 = 0.071$), 62.4% of the variance in continuance intention ($R^2 = 0.624$), and 55.2% of the variance in loyalty ($R^2 = 0.552$). The relatively low explained variance in attitude indicates that AI anxiety and perceived risk account for only a limited psychological-friction component of attitude formation.

Table 6. Structural model results.

Hypothesized direct effects	β	t-value	p-value	95% CI LL-UL	f^2	VIF	Results
H1: A $\rightarrow$ ATT	-0.274	11.191	< .001	[-0.313, -0.214]	0.062	1.293	Supported
H2a/H2b: PR $\rightarrow$ ATT	0.242	6.103	< .001	[0.163, 0.311]	0.050	1.293	H2a not supported; H2b supported
H3: ATT $\rightarrow$ CI	0.557	16.493	< .001	[0.488, 0.619]	0.465	1.768	Supported
H4: SN $\rightarrow$ CI	0.127	5.373	< .001	[0.082, 0.175]	0.033	1.313	Supported
H5: PBC $\rightarrow$ CI	0.239	7.212	< .001	[0.174, 0.305]	0.099	1.518	Supported
H6: CI $\rightarrow$ L	0.667	29.764	< .001	[0.626, 0.713]	1.170	1.000	Supported
Control paths predicting loyalty							
Gender $\rightarrow$ L	0.108	2.760	0.006	[0.030, 0.184]	0.006	1.037	Significant
Age group $\rightarrow$ L	-0.020	0.914	0.361	[-0.064, 0.023]	0.001	1.354	Not significant
Education level $\rightarrow$ L	0.001	0.034	0.973	[-0.042, 0.044]	0.000	1.172	Not significant
Occupational status $\rightarrow$ L	0.016	0.795	0.427	[-0.023, 0.055]	0.000	1.236	Not significant
Usage frequency $\rightarrow$ L	-0.109	4.455	< .001	[-0.158, -0.062]	0.017	1.583	Significant
Membership status $\rightarrow$ L	-0.087	1.674	0.094	[-0.189, 0.015]	0.002	1.181	Not significant
Model fit summary	SRMR = 0.045, d_ULS = 0.960, d_G = 0.449, NFI = 0.833 R ² (L) = 0.552, R ² (CI) = 0.624, R ² (ATT) = 0.071						

Note: A = AI anxiety; ATT = attitude; PR = perceived risk; PBC = perceived behavioral control; SN = subjective norms; CI = continuance intention; L = loyalty. The perceived risk $\rightarrow$ attitude path was specified as a pair of competing hypotheses.

Attitude was positively associated with continuance intention ($\beta = 0.557$, $p < 0.001$, 95% CI [0.488, 0.619], $f^2 = 0.465$), supporting H3. Continuance intention was positively associated with loyalty ($\beta = 0.667$, $p < 0.001$, 95% CI [0.626, 0.713], $f^2 = 1.170$), supporting H6. Both paths in the proposed post-adoption sequence were therefore supported. Effect magnitudes (f^2) are reported here and interpreted in relation to conventional benchmarks in the Discussion (Cohen, 2013; Hair & Alamer, 2022).

At the attitudinal stage, AI anxiety was negatively associated with attitude ($\beta = -0.274$, $p < 0.001$, 95% CI [-0.313, -0.214], $f^2 = 0.062$), supporting H1. The perceived risk-attitude path was specified as a set of competing directional expectations. Perceived risk was positively associated with attitude ($\beta = 0.242$, $p < 0.001$, 95% CI [0.163, 0.311], $f^2 = 0.050$). Thus, the positive direction specified by the contextual salience expectation was supported (H2b), whereas the deterrence expectation was not supported (H2a). This result identifies the supported direction of the association but does not establish contextual salience as the process underlying it. Given the low explained variance in attitude, these direct associations should be interpreted as statistically significant but substantively bounded components of attitude formation rather than as a comprehensive explanation of users' attitudes toward GenAI. Subjective norms ($\beta = 0.127$, $p < 0.001$, 95% CI [0.082, 0.175], $f^2 = 0.033$) and perceived behavioral control ($\beta = 0.239$, $p < 0.001$, 95% CI [0.174, 0.305], $f^2 = 0.099$) were positively associated with continuance intention, supporting H4 and H5. All inner VIF values remained below conservative thresholds, providing no indication of problematic multicollinearity.

Global model fit indices are reported as supplementary descriptive information rather than as primary evidence of model adequacy, given the contested role of global fit measures in PLS-SEM. The SRMR value was 0.045, below the commonly referenced 0.08 threshold, whereas NFI was 0.833 and should therefore be interpreted cautiously. Additional fit-related indices are reported in Table 6 (d_ULS = 0.960, d_G = 0.449; Henseler et al., 2015). The primary evaluation of the model is therefore based on measurement quality, structural paths, explanatory power, and predictive assessment.

Among the control variables included in the loyalty equation, education level was not significantly associated with loyalty ($\beta = 0.001$, $p = 0.973$, 95% CI [-0.042, 0.044]). Age group ($\beta = -0.020$, $p = 0.361$), occupational status ($\beta = 0.016$, $p = 0.427$), and membership status ($\beta = -0.087$, $p = 0.094$) were also not significantly associated with loyalty. By contrast, gender ($\beta = 0.108$, $p = 0.006$, 95% CI [0.030, 0.184]) and usage frequency ($\beta = -0.109$, $p < 0.001$, 95% CI [-0.158, -0.062]) showed statistically significant associations with loyalty.

4.3. Indirect effect analysis

Indirect associations were estimated to examine the pathways specified in the structural model (Hair et al., 2017; Preacher & Hayes, 2008). Table 7 reports the estimated indirect and sequential effects.

The estimated indirect association between AI anxiety and continuance intention through attitude was negative ($\beta = -0.152$, $p < 0.001$, 95% CI [-0.183, -0.116]). The corresponding association for

Table 7. Indirect effects.

Paths	β	t-value	p-value	95% CI LL-UL
A $\rightarrow$ ATT $\rightarrow$ CI	-0.152	8.852	< .001	[-0.183, -0.116]
ATT $\rightarrow$ CI $\rightarrow$ L	0.408	14.702	< .001	[0.353, 0.462]
PBC $\rightarrow$ CI $\rightarrow$ L	0.174	7.161	< .001	[0.128, 0.224]
PR $\rightarrow$ ATT $\rightarrow$ CI	0.136	5.942	< .001	[0.091, 0.178]
SN $\rightarrow$ CI $\rightarrow$ L	0.094	5.344	< .001	[0.059, 0.129]
A $\rightarrow$ ATT $\rightarrow$ CI $\rightarrow$ L	-0.111	8.342	< .001	[-0.136, -0.084]
PR $\rightarrow$ ATT $\rightarrow$ CI $\rightarrow$ L	0.100	5.800	< .001	[0.067, 0.131]

Note: A = AI anxiety; ATT = attitude; PR = perceived risk; PBC = perceived behavioral control; SN = subjective norms; CI = continuance intention; L = loyalty. Confidence intervals are based on bootstrapping with 10,000 resamples. Given the cross-sectional design, indirect effects are interpreted as model-consistent indirect associations rather than evidence of causal mediation.

Table 8. Cross-validated predictive ability test (CVPAT).

Constructs	PLS loss	IA loss	Average loss difference	t-value	p-value
PLS-SEM vs indicator average (IA)					
Attitude (ATT)	0.759	0.798	-0.039	4.185	< .001
Continuance intention (CI)	0.670	0.945	-0.275	14.487	< .001
Loyalty (L)	0.799	0.969	-0.170	12.464	< .001
Overall	0.743	0.904	-0.161	13.416	< .001
PLS-SEM vs linear model (LM)					
Attitude (ATT)	0.510	0.759	-0.249	10.522	< .001
Continuance intention (CI)	0.569	0.670	-0.101	6.913	< .001
Loyalty (L)	0.653	0.799	-0.147	8.211	< .001
Overall	0.577	0.743	-0.166	10.739	< .001

Note: IA = Indicator average; LM = Linear model; CVPAT with 10 folds and 10 repetitions.

perceived risk was positive ($\beta = 0.136$, $p < 0.001$, 95% CI [0.091, 0.178]). This indirect association follows the same direction as the supported direct path (H2b). It does not, however, establish contextual salience as the process underlying either association.

The estimated indirect associations with loyalty through continuance intention were positive for attitude ($\beta = 0.408$, $p < 0.001$, 95% CI [0.353, 0.462]), perceived behavioral control ($\beta = 0.174$, $p < 0.001$, 95% CI [0.128, 0.224]), and subjective norms ($\beta = 0.094$, $p < 0.001$, 95% CI [0.059, 0.129]). The sequential indirect association between AI anxiety and loyalty through attitude and continuance intention was negative ($\beta = -0.111$, $p < 0.001$, 95% CI [-0.136, -0.084]). The corresponding sequential indirect association for perceived risk was positive ($\beta = 0.100$, $p < 0.001$, 95% CI [0.067, 0.131]). These estimates are consistent with the pathways specified in the model, although the cross-sectional design does not permit them to be interpreted as evidence of causal mediation.

4.4. Predictive model assessment

Explanatory and predictive performance are theoretically distinct properties of a structural model, and conflating them risks overstating generalizability (Shmueli et al., 2016, 2019; Shmueli & Koppius, 2011). Whereas R^2 values indicate in-sample explanatory power, out-of-sample predictive performance requires separate assessment (Hair et al., 2019; Hair & Alamer, 2022). We therefore employed two complementary procedures, CVPAT and PLSpredict, both using ten-fold cross-validation with ten repetitions (Sharma et al., 2023; Shmueli et al., 2019). Both procedures benchmark model performance against alternative specifications.

4.5. Cross-validated performance test (CVPAT)

The structural model was evaluated against the indicator average (IA) and linear model (LM) benchmarks (Liengard et al., 2021). As shown in Table 8, the PLS-SEM model outperformed the IA benchmark across all endogenous constructs (overall average loss difference = -0.161, $t = 13.416$, $p < 0.001$), with the largest improvements observed for continuance intention (average loss difference = -0.275) and loyalty (average loss difference = -0.170). Against the more demanding LM benchmark, the model again outperformed the comparison specification across all constructs (overall average loss

Table 9. Predictive relevance across indicators.

Indicators	Q ² predict	RMSE		Difference (PLS-LM)
		PLS-SEM	LM	
ATT1	0.043	0.712	0.854	−0.161
ATT2	0.054	0.718	0.883	−0.165
ATT3	0.048	0.711	0.875	−0.164
CI1	0.308	0.714	0.782	−0.068
CI2	0.285	0.749	0.800	−0.051
CI3	0.283	0.797	0.871	−0.074
L2	0.210	0.738	0.816	−0.078
L3	0.151	0.874	0.972	−0.098
L4	0.173	0.806	0.887	−0.081

Note: ATT = attitude; CI = continuance intention; L = loyalty; LM = linear model; Q²predict = Stone–Geisser's predictive relevance statistic; RMSE = root mean squared error. Predictive relevance for loyalty is reported only for the retained loyalty indicators used in the final measurement model.

difference = −0.166, $t = 10.739$, $p < 0.001$), indicating predictive value beyond both the indicator-average and linear benchmarks (Sharma et al., 2023).

4.6. PLSpredict analysis

Item-level predictive performance was assessed following Shmueli et al. (2019), with Q²predict values derived from holdout-sample cross-validation (Geisser, 1974; Stone, 1974). RMSE values were also benchmarked against the linear model. As reported in Table 9, all Q²predict values were positive, indicating predictive relevance across the nine retained endogenous indicators (Hair et al., 2019). Predictive relevance was strongest for the continuance intention indicators (Q²predict = 0.283–0.308), followed by the retained loyalty indicators L2–L4 (Q²predict = 0.151–0.210), whereas the attitude indicators showed more modest values (Q²predict = 0.043–0.054). RMSE differences were negative across all nine retained endogenous indicators, ranging from −0.051 to −0.165, indicating lower prediction errors for the PLS-SEM model than for the linear benchmark at the indicator level. Both diagnostics therefore point to out-of-sample predictive performance across the endogenous constructs and their retained indicators.

5. Discussion

Attitude showed the strongest association with continuance intention among the three TPB predictors. Respondents' overall evaluations of the focal platform were more strongly associated with continuance intention than either perceived social pressure or perceived behavioral control. Using Cohen's (2013) benchmarks, the effects of AI anxiety ($f^2 = 0.062$) and perceived risk ($f^2 = 0.050$) on attitude, and of subjective norms ($f^2 = 0.033$) and perceived behavioral control ($f^2 = 0.099$) on continuance intention, are small. Although statistically significant, these effects are substantively modest within their respective equations. Accordingly, the theoretical contribution of these relationships lies primarily in their direction and specification, particularly the positive perceived risk–attitude association, rather than in the magnitude of their effects. Consistent with this modest effect magnitude, AI anxiety and perceived risk explained 7.1% of the variance in attitude ($R^2 = 0.071$), indicating that the model captures a limited psychological-friction component rather than providing a comprehensive account of attitude toward GenAI. The present data cannot identify which omitted factors account for the remaining variance or establish whether unmeasured positive determinants are more influential than the psychological frictions examined here. AI anxiety and perceived risk should therefore be interpreted as accounting for a bounded, domain-specific component of attitude within the extended TPB model.

Continuance intention was strongly associated with loyalty ($\beta = 0.667$), consistent with the ordering specified in the model. As Bhattacharjee (2001) argues, continuance intention captures the propensity to keep using a system. However, that propensity does not, by itself, distinguish between users who return out of genuine preference and those who return simply because no clearly superior alternative has displaced the focal option. Conceptually, loyalty can encompass greater stability and relative preference than continuance intention (Dick & Basu, 1994; Oliver, 1999). The retained loyalty construct, however, is limited to affective attachment and personal significance and does not capture behavioral

or cognitive forms of loyalty. Although prior loyalty research associates repeated choice with habit, comparative preference, and resistance to switching (Jones et al., 2000; Limayem et al., 2007), these outcomes were not measured in the present study. The observed coefficient therefore indicates covariance between continuance intention and the retained form of loyalty, not a progression from continued use to comparative preference or switching resistance. Given that both constructs were measured through the same self-report survey at a single point in time, the magnitude of this association should also be interpreted with caution because same-source measurement may have contributed to the observed relationship.

AI anxiety was negatively associated with attitude ($\beta = -0.274$), in line with the model's theoretical expectation. This association is theoretically plausible in AI-assisted travel planning, where users evaluate probabilistic and sometimes opaque outputs in decisions involving time, cost, and experience (Burrell, 2016; Glikson & Woolley, 2020; Weidinger et al., 2022). Prior research similarly associates AI anxiety with less favorable responses, lower acceptance, reduced willingness to engage, and discontinuance-oriented responses across AI-adoption and GenAI contexts (Cengiz & Peker, 2025; Wang & Wang, 2022; Zhang & Tong, 2025). The present result extends this evidence to users' attitudes toward a focal GenAI platform.

The perceived-risk finding warrants particular theoretical attention. Because the perceived risk–attitude path was specified as a set of competing directional expectations, the positive direction of the association ($\beta = 0.242$) supports the contextual salience expectation (H2b) rather than the deterrence expectation (H2a). Classic consumer-risk research conceptualizes perceived risk as an appraisal of uncertainty and possible adverse consequences and associates it with information search and more extensive evaluation (Bauer, 1967; Mitchell, 1999). Privacy-calculus research similarly examines privacy concerns alongside other considerations associated with willingness to transact online (Dinev & Hart, 2006). More broadly, research on algorithm aversion and algorithm appreciation shows that evaluative responses to algorithmic systems vary across tasks, users, and decision contexts rather than following a uniformly negative pattern (Castelo et al., 2019; Dietvorst et al., 2015; Logg et al., 2019). This pattern is particularly relevant in AI-assisted travel planning, where inaccurate, incomplete, or misleading outputs may carry financial, temporal, and experiential consequences. Perceived risk and attitude nevertheless represent different evaluations: perceived risk concerns possible adverse outcomes, whereas attitude reflects the user's overall evaluation of using the platform.

The positive coefficient therefore indicates that higher self-reported perceived risk covaried with more favorable attitudes in this sample; it does not show that perceived risk improved attitude or that contextual salience produced the association. One possible account is that the subjective salience of AI-assisted travel planning varies across users, such that recognition of the measured risks may coexist with favorable attitudes based on other considerations. Because subjective contextual salience was not measured, this account cannot be identified as the process underlying the association. The model also contains no measures of perceived benefit, usefulness, perceived competence, objective risk knowledge, or risk calibration. Limited risk awareness does not straightforwardly explain the positive covariance because lower awareness would ordinarily be reflected in lower perceived-risk scores. Nevertheless, objective risk knowledge was not assessed, and the active-user screening criteria neither demonstrate risk literacy nor exclude early-adopter characteristics. Perceived competence also remains a directionally indeterminate omitted variable, and its role cannot be established from the present data.

5.1. Theoretical implications

The findings offer three major theoretical implications for post-adoption research in GenAI. First, they qualify the conventional deterrence treatment of perceived risk in e-service adoption research, where perceived risk is generally associated with less favorable adoption-related evaluations and intentions (Featherman & Pavlou, 2003; Kim et al., 2008; Pavlou & Fygenson, 2006). Recent studies of AI-related worries likewise document concerns and uncertainty surrounding AI use, although they do not test the perceived risk–attitude relationship examined here (Ayed et al., 2025; Ayed et al., 2025; Batran et al., 2025). By specifying deterrence and contextual salience as competing directional expectations, the present study shows that the observed positive association is consistent with the direction specified by the

contextual salience expectation rather than the deterrence expectation. It does not establish contextual salience or any other unmeasured process as the explanation for that association. The theoretical implication is therefore not that perceived risk should be reclassified as a uniformly positive antecedent of attitude, but that uniformly negative specifications may not describe every GenAI use setting. This result motivates future research that directly measures and tests the roles of contextual salience, perceived benefit, usefulness, competence, objective risk knowledge, experience, and switching conditions in the risk–attitude relationship.

Second, the findings refine the post-adoption application of TPB. The framework has been applied extensively to adoption and continuance intention, and its core constructs remain relevant in the present context. The model extends this structure by incorporating AI anxiety and perceived risk as distinct antecedents of attitude and by examining loyalty as a post-adoption construct associated with continuance intention. For technologies characterized by opacity, probabilistic outputs, and evaluative uncertainty, this extension identifies a limited psychological-friction component that is not represented explicitly in the standard TPB specification (Ajzen, 2020; Conner & Armitage, 1998). The contrasting associations of AI anxiety and perceived risk with attitude further show why these affective and cognitive constructs should not be treated as interchangeable. At the same time, their small effect sizes and the low explained variance in attitude limit any claim that they provide a comprehensive extension of attitude formation in GenAI contexts.

Finally, the findings reinforce the need to distinguish continuance intention from loyalty as a conceptually broader post-adoption construct. Conceptually, loyalty may encompass affective and preferential commitment (Dick & Basu, 1994; Oliver, 1999). Operationally, however, the final construct examined here captures emotional attachment and personal significance toward the focal platform and does not measure comparative preference, resistance to switching, habitual re-engagement, or observed platform choice. The measurement and discriminant-validity results indicate that this retained form of loyalty is empirically distinguishable from continuance intention. The strong cross-sectional association between the constructs does not demonstrate that continuance intention develops into loyalty over time. The contribution therefore lies in extending the post-adoption model beyond continuance intention while maintaining a bounded interpretation of the specific form of loyalty measured in the study.

5.2. Practical implications

The findings suggest that providers should look beyond adoption and usage counts to users' evaluations of focal GenAI services. Because attitude had the strongest association with continuance intention among the TPB predictors, design and service practices that support informed and favorable evaluations warrant consideration. Possible measures include structured onboarding, guided first-use scenarios, explicit expectation-setting, output transparency, and clear error communication. The present study did not test these interventions, and their effects on attitude, continuance intention, and loyalty should therefore be evaluated directly before they are treated as effective retention strategies.

The negative association between AI anxiety and attitude suggests a possible role for support practices that help users manage uncertainty during interaction. Potential measures include realistic explanations of GenAI's capabilities and limitations, guided prompting support, and accessible output-verification procedures. Interface features such as prompt templates, feedback mechanisms, and verification checklists may also help users interpret outputs, but the present data do not show that these interventions reduce anxiety or improve evaluations. Their effectiveness should therefore be examined through experimental or longitudinal evaluation rather than inferred from the observed association. In organizational and educational settings, responsible-use guidance and competence-building initiatives may serve as useful safeguards while their effects on evaluative and post-adoption outcomes are tested.

The positive perceived risk–attitude association should not be interpreted as a reason to increase users' perceptions of risk. It indicates that users reporting greater privacy-, data-control-, uncertainty-, and output-related concerns also reported more favorable attitudes in this sample. From a responsible-design perspective, providers should communicate system uncertainty, data practices, and output limitations clearly rather than obscure them. Practices such as source attribution, uncertainty signaling, and accessible explanations of data use may support informed assessment, but the present study does not

show that such practices improve attitude, calibrate risk perceptions, or increase loyalty. These outcomes require direct evaluation.

The control-variable results should be treated as auxiliary rather than as focal theoretical findings. The model did not detect a statistically significant association between membership status and the measured form of loyalty. One tentative interpretation is that paying for access may reflect functional or task-specific reasons for using a GenAI service without necessarily implying stronger affective attachment or personal significance toward the platform. In a setting where users can access free tiers and potentially use multiple platforms with overlapping functionality, subscription status alone may therefore be an imperfect indicator of the form of loyalty examined here. Explanations involving free-tier availability, functional similarity across platforms, or task-driven subscription decisions were not tested in the present study and should therefore be examined directly before this interpretation can be substantiated. The significant control paths involving gender and usage frequency should likewise be treated as indications warranting further investigation rather than as a basis for subgroup or retention strategies. Research designed specifically to examine user heterogeneity could determine whether and why these associations recur across samples and service settings.

5.3. Limitations and future research

Several limitations should be considered when interpreting these findings, and each points toward a productive research direction. The cross-sectional design measures loyalty and its hypothesized antecedents at a single point in time and therefore cannot establish the ordering specified in the model. Loyalty is often conceptualized as developing through repeated evaluative episodes over time; a cross-sectional design therefore provides a static snapshot of a process that may unfold dynamically (Rindfleisch et al., 2008). Longitudinal designs that track the same users across multiple time points would allow possible progression through Oliver's (1999) loyalty stages, including the action stage not measured here, to be examined more directly rather than inferred from cross-sectional associations. Such designs would also clarify whether the prominence of attitude documented here persists as users accumulate experience or attenuates as normative consensus around GenAI use stabilizes. The model also explained only a modest share of the variance in attitude, indicating that AI anxiety and perceived risk capture the psychological-friction component of attitude formation rather than its full set of antecedents. Subsequent studies should incorporate additional attitude-related predictors, such as perceived usefulness, trust, output quality, enjoyment, AI literacy, prior experience, and expectation confirmation.

The sample is confined to Generation Z users in Turkey, a theoretically motivated choice that nonetheless limits generalizability in ways that matter for the study's most consequential finding. The positive perceived risk-attitude relationship documented here supports the contextual salience expectation rather than the deterrence expectation, but the present design cannot explain why this pattern emerges in this GenAI use context. One plausible explanation is that this sample combines relatively high familiarity with digital technologies with low switching barriers across alternative platforms (Burnham et al., 2003; Pavlou et al., 2007). Because perceived benefit, user experience, perceived competence, risk calibration, and switching cost were not directly measured, these factors should be treated as candidate boundary conditions rather than confirmed explanatory factors. Where perceived benefit is lower, experience is more limited, or switching costs are higher, the standard negative risk-attitude relationship may reassert itself. An important next step is therefore to test these conditions explicitly and examine whether the contextual salience expectation generalizes to older cohorts, lower-familiarity populations, or higher-stakes deployment contexts. Such research would clarify when perceived risk functions as a deterrent and when it operates as a salience-related evaluative signal that coexists with favorable evaluation.

The final loyalty construct captures affective attachment and the personal significance of the focal GenAI platform. It does not directly measure comparative preference, resistance to alternative platforms, cognitive loyalty, behavioral loyalty, habitual re-engagement, or observed platform choice. The findings should therefore not be generalized to these unmeasured forms of loyalty. Future research could combine affective indicators with measures of comparative preference and resistance to alternatives, as well as behavioral indicators such as return frequency, platform exclusivity, and observed

platform choice. Longitudinal and behavioral data would also help determine whether affective attachment and personal significance are reflected in subsequent platform choices over time.

A final limitation concerns CMB. Although procedural remedies and statistical diagnostics were used to mitigate concerns about CMB, the data were collected via a single self-report survey at a single time point. Same-source bias therefore cannot be fully ruled out, particularly when interpreting the strong associations among attitude, continuance intention, and loyalty. The cross-sectional design and self-reported nature of the data mean that causal inferences should not be drawn. In particular, the observed strong path coefficients, especially between continuance intention and loyalty, could be inflated by CMB. Future studies should use multi-source, longitudinal, or experimental designs to separate predictor and outcome measurement, validate self-reported loyalty against behavioral data, and provide stronger evidence regarding the ordering specified in the post-adoption model.

6. Conclusion

Loyalty to generative AI does not automatically follow from access, novelty, or even repeated use. The cross-sectional findings show that attitude was positively associated with continuance intention, which was in turn positively associated with the retained loyalty construct, capturing affective attachment and personal significance. This pattern is consistent with the hypothesized associations among attitude, continuance intention, and loyalty but does not establish their temporal ordering.

AI anxiety was negatively associated with attitude and had a negative sequential indirect association with loyalty through attitude and continuance intention. Perceived risk, contrary to the deterrence expectation, was positively associated with attitude and had a positive sequential indirect association through the same constructs. The positive perceived risk^{**}–^{**}attitude coefficient supports the direction specified by the contextual salience expectation but does not identify contextual salience or another unmeasured factor as the process underlying the association. Both AI anxiety and perceived risk had small effects on attitude, and together they explained only a limited share of its variance.

The contribution of the study therefore lies in distinguishing continuance intention from affective attachment and personal significance, specifying competing directional expectations for perceived risk, and documenting the contrasting associations of AI anxiety and perceived risk with attitude. Longitudinal, multi-source, and behavioral designs are needed to test temporal ordering, examine the unmeasured explanations for the perceived-risk result, and determine whether the measured form of loyalty is reflected in subsequent platform choice.

Ethical considerations

This study was based on an anonymous, voluntary survey of adult participants and did not involve medical procedures, interventions, vulnerable populations, or the collection of personally identifiable information. In accordance with the written policy of the *Ibn Haldun University Scientific Research and Publication Ethics Committee for Social Sciences and Humanities*, formal ethics committee approval was not required because the study involved only anonymous, non-sensitive survey data collected from voluntary adult participants, without intervention or the collection of personally identifiable information. Accordingly, the study was exempt from formal ethics committee review, and no ethics approval reference number was issued. Before completing the questionnaire, respondents were informed about the academic purpose of the study, and electronic informed consent was obtained. All responses were treated confidentially and used solely for research purposes. The study was conducted in accordance with institutional ethical principles and internationally accepted standards for research involving human participants.

Author contributions

CRedit: **Ibrahim Arpacı**: Conceptualization, Formal analysis, Project administration, Supervision, Writing – original draft; **Azize Sahin**: Conceptualization, Formal analysis, Methodology, Resources, Software, Writing – original draft; **Ekrem Tatoglu**: Conceptualization, Investigation, Methodology, Project administration, Supervision, Writing – review & editing; **Aysun Sahin**: Conceptualization, Formal analysis, Methodology, Resources, Writing – original draft.

AI declaration

During the preparation of this manuscript, ChatGPT and Grammarly Pro were used only for language editing, grammar checking, and readability improvement. These tools were not used to generate research data, conduct statistical analyses, interpret results, or produce original scholarly claims. The authors reviewed and approved all AI-assisted edits and take full responsibility for the final content of the manuscript.

Disclosure statement

No potential conflict of interest was reported by the author(s).

ORCID

Ibrahim Arpacı  <http://orcid.org/0000-0001-6513-4569>

Azize Sahin  <http://orcid.org/0000-0002-3115-6812>

Ekrem Tatoglu  <http://orcid.org/0000-0002-9119-3252>

Aysun Sahin  <http://orcid.org/0000-0003-3285-1232>

Data availability statement

The de-identified dataset supporting the findings of this study is publicly available on OSF at <https://osf.io/kvr5z/>.

References

- Ajzen, I. (1991). The theory of planned behavior. *Organizational Behavior and Human Decision Processes*, 50(2), 179–211. [https://doi.org/10.1016/0749-5978\(91\)90020-T](https://doi.org/10.1016/0749-5978(91)90020-T)
- Ajzen, I. (2001). Nature and operation of attitudes. *Annual Review of Psychology*, 52(1), 27–58. <https://doi.org/10.1146/annurev.psych.52.1.27>
- Ajzen, I. (2002). Perceived behavioral control, self-efficacy, locus of control, and the theory of planned behavior. *Journal of Applied Social Psychology*, 32(4), 665–683. <https://doi.org/10.1111/j.1559-1816.2002.tb00236.x>
- Ajzen, I. (2020). The theory of planned behavior: Frequently asked questions. *Human Behavior and Emerging Technologies*, 2(4), 314–324. <https://doi.org/10.1002/hbe2.195>
- Al Amin, M., Muzareba, A. M., Chowdhury, I. U., & Khondkar, M. (2024). Understanding e-satisfaction, continuance intention, and e-loyalty toward mobile payment application during COVID-19: An investigation using the electronic technology continuance model. *Journal of Financial Services Marketing*, 29(2), 318–340. <https://doi.org/10.1057/s41264-022-00197-2>
- Al-Emran, M., Arpacı, I., & Salloum, S. A. (2020). An empirical examination of continuous intention to use m-learning: An integrated model. *Education and Information Technologies*, 25(4), 2899–2918. <https://doi.org/10.1007/s10639-019-10094-2>
- Al-Sharafi, M. A., Al-Emran, M., Iranmanesh, M., Al-Qaysi, N., Iahad, N. A., & Arpacı, I. (2023). Understanding the impact of knowledge management factors on the sustainable use of AI-based chatbots for educational purposes using a hybrid SEM-ANN approach. *Interactive Learning Environments*, 31(10), 7491–7510. <https://doi.org/10.1080/10494820.2022.2075014>
- Amer, B., Ayed, A., Malak, M., & Bashtawy, M. (2025). Nursing informatics competency and self-efficacy in clinical practice among nurses in Palestinian hospitals. *Hospital Topics*, 103(2), 122–129. <https://doi.org/10.1080/00185868.2023.2252974>
- Appel, G., Grewal, L., Hadi, R., & Stephen, A. T. (2020). The future of social media in marketing. *Journal of the Academy of Marketing Science*, 48(1), 79–95. <https://doi.org/10.1007/s11747-019-00695-1>
- Armstrong, J. S., & Overton, T. S. (1977). Estimating nonresponse bias in mail surveys. *Journal of Marketing Research*, 14(3), 396–402. <https://doi.org/10.1177/002224377701400320>
- Ashfaq, M., Yun, J., Yu, S., & Loureiro, S. M. C. (2020). I, chatbot: Modeling the determinants of users' satisfaction and continuance intention of AI-powered service agents. *Telematics and Informatics*, 54, 101473. <https://doi.org/10.1016/j.tele.2020.101473>
- Ayed, A., Batran, A., Aqtam, I., Malak, M. Z., Ejheisheh, M. A., Farajallah, M., Farraj, L., & Alkhatib, S. (2025). Perceived worries in the adoption of artificial intelligence among nurses in neonatal intensive care units. *BMC Nursing*, 24(1), 777. <https://doi.org/10.1186/s12912-025-03318-z>
- Ayed, A., Ejheisheh, M. A., Al-Amer, R., Aqtam, I., Ali, A. M., Othman, E. H., Farajallah, M., Qaddumi, J., & Batran, A. (2025). Insights into the relationship between anxiety and attitudes toward artificial intelligence among nursing students. *BMC Nursing*, 24(1), 812. <https://doi.org/10.1186/s12912-025-03490-2>

- Back, K. J., & Parks, S. C. (2003). A brand loyalty model involving cognitive, affective, and conative brand loyalty and customer satisfaction. *Journal of Hospitality & Tourism Research*, 27(4), 419–435. <https://doi.org/10.1177/10963480030274003>
- Bandura, A. (1986). *Social foundations of thought and action: A social cognitive theory*. Prentice-Hall.
- Batran, A., Al-Humran, S. M., Malak, M. Z., & Ayed, A. (2022). The relationship between nursing informatics competency and clinical decision-making among nurses in West Bank, Palestine. *Computers, Informatics, Nursing: CIN*, 40(8), 547–553. <https://doi.org/10.1097/CIN.0000000000000890>
- Batran, A., Ayed, A., Aqtam, I., Al-Amer, R., Othman, E. H., Ejheisheh, M. A., Al Daamsa, H. N., Alkhatib, S., & Farajallah, M. (2025). Insights into perceived worries regarding the adoption of artificial intelligence among intensive care unit nurses in the West Bank. *SAGE Open Nursing*, 11, 23779608251376177. <https://doi.org/10.1177/23779608251376177>
- Bauer, R. A. (1967). Consumer behavior as risk taking. In M. J. Baker (Ed.), *Marketing: Critical perspectives on business and management* (Vol. 3, pp. 13–21). Routledge.
- Bhattacharjee, A. (2001). An empirical analysis of the antecedents of electronic commerce service continuance. *Decision Support Systems*, 32(2), 201–214. [https://doi.org/10.1016/S0167-9236\(01\)00111-7](https://doi.org/10.1016/S0167-9236(01)00111-7)
- Blut, M., Wang, C., Wunderlich, N. V., & Brock, C. (2021). Understanding anthropomorphism in service provision: A meta-analysis of physical robots, chatbots, and other AI. *Journal of the Academy of Marketing Science*, 49(4), 632–658. <https://doi.org/10.1007/s11747-020-00762-y>
- Boulus-Rødje, N., Cranefield, J., Doyle, C., & Fleron, B. (2024). GenAI and me: The hidden work of building and maintaining an augmentative partnership. *Personal and Ubiquitous Computing*, 28(6), 861–874. <https://doi.org/10.1007/s00779-024-01810-y>
- Brislin, R. W. (1986). The wording and translation of research instruments. In W. J. Lonner & J. W. Berry (Eds.), *Field methods in cross-cultural research* (pp. 137–164). Sage Publications.
- Burnham, T. A., Frels, J. K., & Mahajan, V. (2003). Consumer switching costs: A typology, antecedents, and consequences. *Journal of the Academy of Marketing Science*, 31(2), 109–126. <https://doi.org/10.1177/0092070302250897>
- Burrell, J. (2016). How the machine ‘thinks’: Understanding opacity in machine learning algorithms. *Big Data & Society*, 3(1), 1–12. <https://doi.org/10.1177/2053951715622512>
- Calder, B. J., Phillips, L. W., & Tybout, A. M. (1981). Designing research for application. *Journal of Consumer Research*, 8(2), 197–207. <https://doi.org/10.1086/208856>
- Castelo, N., Bos, M. W., & Lehmann, D. R. (2019). Task-dependent algorithm aversion. *Journal of Marketing Research*, 56(5), 809–825. <https://doi.org/10.1177/0022243719851788>
- Cengiz, S., & Peker, A. (2025). Generative artificial intelligence acceptance and artificial intelligence anxiety among university students: The sequential mediating role of attitudes toward artificial intelligence and literacy. *Current Psychology*, 44(9), 7991–8000. <https://doi.org/10.1007/s12144-025-07433-7>
- Chen, J., He, M., & Sun, J. (2025). AI anxiety and knowledge payment: The roles of perceived value and self-efficacy. *BMC Psychology*, 13(1), 208. <https://doi.org/10.1186/s40359-025-02510-9>
- Cohen, J. (2013). *Statistical power analysis for the behavioral sciences*. Routledge.
- Compeau, D. R., & Higgins, C. A. (1995). Computer self-efficacy: Development of a measure and initial test. *MIS Quarterly*, 19(2), 189–211. <https://doi.org/10.2307/249688>
- Conner, M., & Armitage, C. J. (1998). Extending the theory of planned behavior: A review and avenues for further research. *Journal of Applied Social Psychology*, 28(15), 1429–1464. <https://doi.org/10.1111/j.1559-1816.1998.tb01685.x>
- Cossío-Silva, F. J., Revilla-Camacho, M. Á., Vega-Vázquez, M., & Palacios-Florencio, B. (2016). Value co-creation and customer loyalty. *Journal of Business Research*, 69(5), 1621–1625. <https://doi.org/10.1016/j.jbusres.2015.10.028>
- Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, 13(3), 319–340. <https://doi.org/10.2307/249008>
- Deloitte. (2024). *Digital consumer trends 2024: Türkiye*. <https://www.deloitte.com/tr/en/Industries/tmt/research/digital-consumer-trends.html>
- Deloitte. (2025). *Gen Z and millennial survey 2025: Living and working with purpose in a transforming world*. <https://www.deloitte.com/global/en/issues/work/content/genzmillennialsurvey.html>
- Dick, A. S., & Basu, K. (1994). Customer loyalty: Toward an integrated conceptual framework. *Journal of the Academy of Marketing Science*, 22(2), 99–113. <https://doi.org/10.1177/0092070394222001>
- Dietvorst, B. J., Simmons, J. P., & Massey, C. (2015). Algorithm aversion: People erroneously avoid algorithms after seeing them err. *Journal of Experimental Psychology: General*, 144(1), 114–126. <https://doi.org/10.1037/xge0000033>
- Dinev, T., & Hart, P. (2006). An extended privacy calculus model for e-commerce transactions. *Information Systems Research*, 17(1), 61–80. <https://doi.org/10.1287/isre.1060.0080>
- Douglas, S. P., & Craig, C. S. (2007). Collaborative and iterative translation: An alternative approach to back translation. *Journal of International Marketing*, 15(1), 30–43. <https://doi.org/10.1509/jimk.15.1.030>

- Dwivedi, Y. K., Hughes, L., Ismagilova, E., Aarts, G., Coombs, C., Crick, T., Duan, Y., Dwivedi, R., Edwards, J., Eirug, A., Galanos, V., Ilavarasan, P. V., Janssen, M., Jones, P., Kar, A. K., Kizgin, H., Kronemann, B., Lal, B., Lucini, B., ... 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, 101994. <https://doi.org/10.1016/j.ijinfomgt.2019.08.002>
- Dwivedi, Y. K., Kshetri, N., Hughes, L., Slade, E. L., Jeyaraj, A., Kar, A. K., Baabdullah, A. M., Koohang, A., Raghavan, V., Ahuja, M., Albanna, H., Albashrawi, M. A., Al-Busaidi, A. S., Balakrishnan, J., Barlette, Y., Basu, S., Bose, I., Brooks, L., Buhalis, D., ... Wright, R. (2023). "So what if ChatGPT wrote it?" Multidisciplinary perspectives on opportunities, challenges and implications of generative conversational AI for research, practice and policy. *International Journal of Information Management*, 71, 102642. <https://doi.org/10.1016/j.ijinfomgt.2023.102642>
- Eagly, A. H., & Chaiken, S. (1993). *The psychology of attitudes*. Harcourt Brace Jovanovich College Publishers.
- Farrell, J., & Klemperer, P. (2007). Coordination and lock-in: Competition with switching costs and network effects. In M. Armstrong & R. Porter (Eds.), *Handbook of industrial organization* (Vol. 3, pp. 1967–2072). Elsevier. [https://doi.org/10.1016/S1573-448X\(06\)03031-7](https://doi.org/10.1016/S1573-448X(06)03031-7)
- Featherman, M. S., & Pavlou, P. A. (2003). Predicting e-services adoption: A perceived risk facets perspective. *International Journal of Human-Computer Studies*, 59(4), 451–474. [https://doi.org/10.1016/S1071-5819\(03\)00111-3](https://doi.org/10.1016/S1071-5819(03)00111-3)
- Festinger, L. (1954). A theory of social comparison processes. *Human Relations*, 7(2), 117–140. <https://doi.org/10.1177/001872675400700202>
- Fidan, Ü. (2024). Assessment of Türkiye's digitalization performance within the framework of the UN Sustainable Development Index. *Uluslararası Yönetim Bilişim Sistemleri ve Bilgisayar Bilimleri Dergisi*, 8(1), 1–14. <https://doi.org/10.33461/uybisbbd.1373965>
- Fornell, C., & Larcker, D. F. (1981). Evaluating structural equation models with unobservable variables and measurement error. *Journal of Marketing Research*, 18(1), 39–50. <https://doi.org/10.1177/002224378101800104>
- Foroughi, B., Vu, H. T. M., & Thaichon, P. (2026). Building trust for sustained generative AI travel adoption. *Tourism Review*, 81(2), 718–754. <https://doi.org/10.1108/TR-05-2025-0554>
- Geisser, S. (1974). A predictive approach to the random effect model. *Biometrika*, 61(1), 101–107. <https://doi.org/10.1093/biomet/61.1.101>
- Gerlich, M. (2024). Public anxieties about AI: Implications for corporate strategy and societal impact. *Administrative Sciences*, 14(11), 288. <https://doi.org/10.3390/admsci14110288>
- Glasman, L. R., & Albarracín, D. (2006). Forming attitudes that predict future behavior: A meta-analysis of the attitude-behavior relation. *Psychological Bulletin*, 132(5), 778–822. <https://doi.org/10.1037/0033-2909.132.5.778>
- Glikson, E., & Woolley, A. W. (2020). Human trust in artificial intelligence: Review of empirical research. *Academy of Management Annals*, 14(2), 627–660. <https://doi.org/10.5465/annals.2018.0057>
- Gustafsson, A., Johnson, M. D., & Roos, I. (2005). The effects of customer satisfaction, relationship commitment dimensions, and triggers on customer retention. *Journal of Marketing*, 69(4), 210–218. <https://doi.org/10.1509/jmkg.2005.69.4.210>
- Hair, J., & Alamer, A. (2022). Partial least squares structural equation modeling (PLS-SEM) in second language and education research: Guidelines using an applied example. *Research Methods in Applied Linguistics*, 1(3), 100027. <https://doi.org/10.1016/j.rmal.2022.100027>
- Hair, J. F., Hult, G. T. M., Ringle, C. M., & Sarstedt, M. (2017). *A primer on partial least squares structural equation modeling (PLS-SEM)*. Sage.
- Hair, J. F., Risher, J. J., Sarstedt, M., & Ringle, C. M. (2019). When to use and how to report the results of PLS-SEM. *European Business Review*, 31(1), 2–24. <https://doi.org/10.1108/EBR-11-2018-0203>
- Han, H., Kim, S., Hailu, T. B., Al-Ansi, A., Lee, J., & Kim, J. J. (2024). Effects of cognitive, affective and normative drivers of artificial intelligence ChatGPT on continuous use intention. *Journal of Hospitality and Tourism Technology*, 15(4), 629–647. <https://doi.org/10.1108/JHTT-11-2023-0363>
- Harman, H. H. (1976). *Modern factor analysis* (3rd ed.). University of Chicago Press.
- Henseler, J., Ringle, C. M., & Sarstedt, M. (2015). A new criterion for assessing discriminant validity in variance-based structural equation modeling. *Journal of the Academy of Marketing Science*, 43(1), 115–135. <https://doi.org/10.1007/s11747-014-0403-8>
- Ivanov, S., Soliman, M., Tuomi, A., Alkathiri, N. A., & Al-Alawi, A. N. (2024). Drivers of generative AI adoption in higher education through the lens of the theory of planned behaviour. *Technology in Society*, 77, 102521. <https://doi.org/10.1016/j.techsoc.2024.102521>
- Jones, M. A., Mothersbaugh, D. L., & Beatty, S. E. (2000). Switching barriers and repurchase intentions in services. *Journal of Retailing*, 76(2), 259–274. [https://doi.org/10.1016/S0022-4359\(00\)00024-5](https://doi.org/10.1016/S0022-4359(00)00024-5)
- Kim, D. J., Ferrin, D. L., & Rao, H. R. (2008). A trust-based consumer decision-making model in electronic commerce: The role of trust, perceived risk, and their antecedents. *Decision Support Systems*, 44(2), 544–564. <https://doi.org/10.1016/j.dss.2007.07.001>

- Kim, M. J., Lee, C. K., & Jung, T. (2020). Exploring consumer behavior in virtual reality tourism using an extended stimulus-organism-response model. *Journal of Travel Research*, 59(1), 69–89. <https://doi.org/10.1177/0047287518818915>
- Kock, N., & Lynn, G. S. (2012). Lateral collinearity and misleading results in variance-based SEM: An illustration and recommendations. *Journal of the Association for Information Systems*, 13(7), 546–580. <https://doi.org/10.17705/1jais.00302>
- Kock, N. (2015). Common method bias in PLS-SEM. *International Journal of e-Collaboration*, 11(4), 1–10. <https://doi.org/10.4018/ijec.2015100101>
- Lalicic, L., & Weismayer, C. (2021). Consumers' reasons and perceived value co-creation of using artificial intelligence-enabled travel service agents. *Journal of Business Research*, 129, 891–901. <https://doi.org/10.1016/j.jbusres.2020.11.005>
- Lee, J. C., Tang, Y., & Jiang, S. (2023). Understanding continuance intention of artificial intelligence (AI)-enabled mobile banking applications: An extension of AI characteristics to an expectation confirmation model. *Humanities and Social Sciences Communications*, 10(1), 333. <https://doi.org/10.1057/s41599-023-01845-1>
- Lee, K. Y., Sheehan, L., Lee, K., & Chang, Y. (2021). The continuation and recommendation intention of artificial intelligence-based voice assistant systems (AIVAS): The influence of personal traits. *Internet Research*, 31(5), 1899–1939. <https://doi.org/10.1108/INTR-06-2020-0327>
- Li, C., Su, J., & Yang, Y. (2026). Understanding AI anxiety based on terror management theory: A meta-analytical construction. *International Journal of Information Management*, 88, 103043. <https://doi.org/10.1016/j.ijinfomgt.2026.103043>
- Li, Y., & Lee, S. O. (2025). Navigating the generative AI travel landscape: The influence of ChatGPT on the evolution from new users to loyal adopters. *International Journal of Contemporary Hospitality Management*, 37(4), 1421–1447. <https://doi.org/10.1108/IJCHM-11-2023-1767>
- Liengaard, B. D., Sharma, P. N., Hult, G. T. M., Jensen, M. B., Sarstedt, M., Hair, J. F., & Ringle, C. M. (2021). Prediction: Coveted, yet forsaken? Introducing a cross-validated predictive ability test in partial least squares path modeling. *Decision Sciences*, 52(2), 362–392. <https://doi.org/10.1111/deci.12445>
- Lim, N. (2003). Consumers' perceived risk: Sources versus consequences. *Electronic Commerce Research and Applications*, 2(3), 216–228. [https://doi.org/10.1016/S1567-4223\(03\)00025-5](https://doi.org/10.1016/S1567-4223(03)00025-5)
- Limayem, M., Hirt, S. G., & Cheung, C. M. K. (2007). How habit limits the predictive power of intention: The case of information systems continuance. *MIS Quarterly*, 31(4), 705–737. <https://doi.org/10.2307/25148817>
- Logg, J. M., Minson, J. A., & Moore, D. A. (2019). Algorithm appreciation: People prefer algorithmic to human judgment. *Organizational Behavior and Human Decision Processes*, 151, 90–103. <https://doi.org/10.1016/j.obhdp.2018.12.005>
- Long, D., & Magerko, B. (2020). *What is AI literacy? Competencies and design considerations* [Paper presentation]. Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems (pp. 1–16), Honolulu, HI, USA. <https://doi.org/10.1145/3313831.3376727>
- McKinsey & Company. (2025). The state of AI: Global survey 2025. <https://www.mckinsey.com/capabilities/quantumblack/our-insights/the-state-of-ai>
- Mercer, A. W., Kreuter, F., Keeter, S., & Stuart, E. A. (2017). Theory and practice in nonprobability surveys. *Public Opinion Quarterly*, 81(S1), 250–271. <https://doi.org/10.1093/poq/nfx024>
- Memon, M. A., Ting, H., Cheah, J. H., Thursamy, R., Chuah, F., & Cham, T. H. (2020). Sample size for survey research: Review and recommendations. *Journal of Applied Structural Equation Modeling*, 4(2), i–xx. [https://doi.org/10.47263/JASEM.4\(2\)01](https://doi.org/10.47263/JASEM.4(2)01)
- Mick, D. G., & Fournier, S. (1998). Paradoxes of technology: Consumer cognizance, emotions, and coping strategies. *Journal of Consumer Research*, 25(2), 123–143. <https://doi.org/10.1086/209531>
- Mishra, A., Shukla, A., Rana, N., Currie, W., & Dwivedi, Y. (2023). Re-examining post-acceptance model of information systems continuance: A revised theoretical model using MASEM approach. *International Journal of Information Management*, 68, 102571. <https://doi.org/10.1016/j.ijinfomgt.2022.102571>
- Mitchell, V. (1999). Consumer perceived risk: Conceptualisations and models. *European Journal of Marketing*, 33(1/2), 163–195. <https://doi.org/10.1108/03090569910249229>
- Mun, I. B., & Hwang, K. H. (2025). Exploring the influence of prompt self-efficacy: Accurate and customized information, perceived ease of use, satisfaction, and continuance intention to use ChatGPT. *International Journal of Human-Computer Interaction*, 41(22), 13952–13963. <https://doi.org/10.1080/10447318.2025.2480823>
- Needleman, S. E. (2026). Gen Z uses AI a lot—even though it freaks them out, a new study shows. *Business Insider*. <https://www.businessinsider.com/ai-gen-z-job-killer-or-opportunity-2026-1>
- Needleman, S. E. (2023). AI could supercharge the Gen Z takeover. *The Wall Street Journal*. <https://www.wsj.com/tech/ai/ai-could-supercharge-gen-z-workers-career-advancement-5c2f6e92>
- Oliver, R. L. (1999). Whence consumer loyalty? *Journal of Marketing*, 63(4_suppl1), 33–44. <https://doi.org/10.1177/00222429990634s105>
- Palinkas, L. A., Horwitz, S. M., Green, C. A., Wisdom, J. P., Duan, N., & Hoagwood, K. (2015). Purposeful sampling for qualitative data collection and analysis in mixed method implementation research. *Administration and Policy in Mental Health*, 42(5), 533–544. <https://doi.org/10.1007/s10488-013-0528-y>

- Pavlou, P. A., & Fygenson, M. (2006). Understanding and predicting electronic commerce adoption: An extension of the theory of planned behavior. *MIS Quarterly*, 30(1), 115–143. <https://doi.org/10.2307/25148720>
- Pavlou, P. A., Liang, H., & Xue, Y. (2007). Understanding and mitigating uncertainty in online exchange relationships: A principal–agent perspective. *MIS Quarterly*, 31(1), 105–136. <https://doi.org/10.2307/25148783>
- Podsakoff, P. M., MacKenzie, S. B., & Podsakoff, N. P. (2012). Sources of method bias in social science research and recommendations on how to control it. *Annual Review of Psychology*, 63(1), 539–569. <https://doi.org/10.1146/annurev-psych-120710-100452>
- Podsakoff, P. M., Podsakoff, N. P., Williams, L. J., Huang, C., & Yang, J. (2024). Common method bias: It's bad, it's complex, it's widespread, and it's not easy to fix. *Annual Review of Organizational Psychology and Organizational Behavior*, 11(1), 17–61. <https://doi.org/10.1146/annurev-orgpsych-110721-040030>
- Popşa, R. E. (2024). Exploring the generation Z travel trends and behavior. *Studies in Business and Economics*, 19(1), 189–189. <https://doi.org/10.2478/sbe-2024-0010>
- Preacher, K. J., & Hayes, A. F. (2008). Asymptotic and resampling strategies for assessing and comparing indirect effects in multiple mediator models. *Behavior Research Methods*, 40(3), 879–891. <https://doi.org/10.3758/BRM.40.3.879>
- Prensky, M. (2001). Digital natives, digital immigrants part 1. *On the Horizon*, 9(5), 1–6. <https://doi.org/10.1108/10748120110424816>
- Rindfleisch, A., Malter, A. J., Ganesan, S., & Moorman, C. (2008). Cross-sectional versus longitudinal survey research: Concepts, findings, and guidelines. *Journal of Marketing Research*, 45(3), 261–279. <https://doi.org/10.1509/jmkr.45.3.261>
- Ringle, C. M., Wende, S., & Becker, J. M. (2024). SmartPLS 4 (Version 4.1.1.5) [Computer software]. SmartPLS GmbH. <https://www.smartpls.com>
- Rochet, J. C., & Tirole, J. (2003). Platform competition in two-sided markets. *Journal of the European Economic Association*, 1(4), 990–1029. <https://doi.org/10.1162/154247603322493212>
- Rogers, E. M. (2003). *Diffusion of innovations* (5th ed.). Free Press.
- Şahin, A., Kitapçı, H., & Zehir, C. (2013). Creating commitment, trust and satisfaction for a brand: What is the role of switching costs in mobile phone market? *Procedia - Social and Behavioral Sciences*, 99, 496–502. <https://doi.org/10.1016/j.sbspro.2013.10.518>
- Sarstedt, M., Ringle, C. M., & Hair, J. F. (2017). Treating unobserved heterogeneity in PLS-SEM: A multi-method approach. In H. Latan & R. Noonan (Eds.), *Partial least squares path modeling* (pp. 197–217). Springer International Publishing. https://doi.org/10.1007/978-3-319-64069-3_9
- Sarstedt, M., Ringle, C. M., & Hair, J. F. (2022). Partial least squares structural equation modeling. In C. Homburg, M. Klarmann, & A. Vomberg (Eds.), *Handbook of market research* (pp. 587–632). Springer International Publishing. https://doi.org/10.1007/978-3-319-57413-4_15
- Saxena, A., & Doleck, T. (2023). A structural model of student continuance intentions in ChatGPT adoption. *Eurasia Journal of Mathematics, Science and Technology Education*, 19(12), em2366. <https://doi.org/10.29333/ejmste/13839>
- Schepman, A., & Rodway, P. (2023). The general attitudes towards artificial intelligence scale (GAAIS): Confirmatory validation and associations with personality, corporate distrust, and general trust. *International Journal of Human–Computer Interaction*, 39(13), 2724–2741. <https://doi.org/10.1080/10447318.2022.2085400>
- Sharma, P. N., Liengaard, B. D., Hair, J. F., Sarstedt, M., & Ringle, C. M. (2023). Predictive model assessment and selection in composite-based modeling using PLS-SEM: Extensions and guidelines for using CVPAT. *European Journal of Marketing*, 57(6), 1662–1677. <https://doi.org/10.1108/EJM-08-2020-0636>
- Shmueli, G., & Koppius, O. R. (2011). Predictive analytics in information systems research. *MIS Quarterly*, 35(3), 553–572. <https://doi.org/10.2307/23042796>
- Shmueli, G., Ray, S., Velasquez Estrada, J. M., & Chatla, S. B. (2016). The elephant in the room: Predictive performance of PLS models. *Journal of Business Research*, 69(10), 4552–4564. <https://doi.org/10.1016/j.jbusres.2016.03.049>
- Shmueli, G., Sarstedt, M., Hair, J. F., Cheah, J. H., Ting, H., Vaithilingam, S., & Ringle, C. M. (2019). Predictive model assessment in PLS-SEM: Guidelines for using PLSpredict. *European Journal of Marketing*, 53(11), 2322–2347. <https://doi.org/10.1108/EJM-02-2019-0189>
- Srinivasan, S. S., Anderson, R., & Ponnnavolu, K. (2002). Customer loyalty in e-commerce: An exploration of its antecedents and consequences. *Journal of Retailing*, 78(1), 41–50. [https://doi.org/10.1016/S0022-4359\(01\)00065-3](https://doi.org/10.1016/S0022-4359(01)00065-3)
- Stone, M. (1974). Cross-validatory choice and assessment of statistical predictions. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 36(2), 111–133. <https://doi.org/10.1111/j.2517-6161.1974.tb00994.x>
- Taylor, S., & Todd, P. A. (1995). Understanding information technology usage: A test of competing models. *Information Systems Research*, 6(2), 144–176. <https://doi.org/10.1287/isre.6.2.144>
- Tong, H., Yu, J., & Shou, M. (2025). Mitigating the effect of AI anxiety on employees' creativity: A social cognitive perspective. *Journal of Digital Management*, 1(1), 7. <https://doi.org/10.1007/s44362-025-00006-5>
- Vatanasombut, B., Igbaria, M., Stylianou, A. C., & Rodgers, W. (2008). Information systems continuance intention of web-based applications customers: The case of online banking. *Information & Management*, 45(7), 419–428. <https://doi.org/10.1016/j.im.2008.03.005>

- Venkatesh, V., Morris, M. G., Davis, G. B., & Davis, F. D. (2003). User acceptance of information technology: Toward a unified view. *MIS Quarterly*, 27(3), 425–478. <https://doi.org/10.2307/30036540>
- Venkatesh, V., Thong, J. Y. L., & Xu, X. (2012). Consumer acceptance and use of information technology: Extending the unified theory of acceptance and use of technology. *MIS Quarterly*, 36(1), 157–178. <https://doi.org/10.2307/41410412>
- Venkatesh, V., & Davis, F. D. (2000). A theoretical extension of the technology acceptance model: Four longitudinal field studies. *Management Science*, 46(2), 186–204. <https://doi.org/10.1287/mnsc.46.2.186.11926>
- Wang, C., Wang, H., Li, Y., Dai, J., Gu, X., & Yu, T. (2025). Factors influencing university students' behavioral intention to use generative artificial intelligence: Integrating the theory of planned behavior and AI literacy. *International Journal of Human–Computer Interaction*, 41(11), 6649–6671. <https://doi.org/10.1080/10447318.2024.2383033>
- Wang, Y. M., Wei, C. L., Lin, H. H., Wang, S. C., & Wang, Y. S. (2022). What drives students' AI learning behavior: A perspective of AI anxiety. *Interactive Learning Environments*, 32(6), 1–17. <https://doi.org/10.1080/10494820.2022.2153147>
- Wang, Y. Y., & Wang, Y. S. (2022). Development and validation of an artificial intelligence anxiety scale: An initial application in predicting motivated learning behavior. *Interactive Learning Environments*, 30(4), 619–634. <https://doi.org/10.1080/10494820.2019.1674887>
- Ward, M. K., & Meade, A. W. (2023). Dealing with careless responding in survey data: Prevention, identification, and recommended best practices. *Annual Review of Psychology*, 74(1), 577–596. <https://doi.org/10.1146/annurev-psych-040422-045007>
- Weidinger, L., Uesato, J., Rauh, M., Griffin, C., Huang, P. S., Mellor, J., Glaese, A., Cheng, M., Balle, B., Kasirzadeh, A., Biles, C., Brown, S., Kenton, Z., Hawkins, W., Stepleton, T., Birhane, A., Hendricks, L. A., Rimell, L., Isaac, W., ... Gabriel, I. (2022). *Taxonomy of risks posed by language models* [Paper presentation]. 2022 ACM Conference on Fairness, Accountability, and Transparency (pp. 214–229), Seoul, Republic of Korea. <https://doi.org/10.1145/3531146.3533088>
- Wessel, M., Adam, M., Benlian, A., Majchrzak, A., & Thies, F. (2025). Generative AI and its transformative value for digital platforms. *Journal of Management Information Systems*, 42(2), 346–369. <https://doi.org/10.1080/07421222.2025.2487315>
- Zhang, T., & Tong, Q. (2025). The technostress of ChatGPT use: How do perceived AI characteristics affect users' discontinuous use through AI anxiety and negative attitudes? *International Journal of Human–Computer Interaction*, 41(16), 9918–9929. <https://doi.org/10.1080/10447318.2024.2429889>
- Zhu, Y., Zhang, J., Wu, J., & Liu, Y. (2022). AI is better when I'm sure: The influence of certainty of needs on consumers' acceptance of AI chatbots. *Journal of Business Research*, 150, 642–652. <https://doi.org/10.1016/j.jbusres.2022.06.044>

About the authors

Ibrahim Arpacı is Professor of Information Systems at Bursa Uludağ University, Turkey, and Research Professor at Korea University, South Korea. His research focuses on information systems, educational technology, cybersecurity, and cyberpsychology. He has published over 100 journal articles and is ranked among the world's top 2% most influential scientists.

Azize Sahin is Associate Professor of Marketing at Istanbul University, Turkey. Her research focuses on brand management, consumer behavior, experiential marketing, and digital marketing. She has published in leading journals and has extensive industry experience in branding, advertising, executive education, and strategic marketing consultancy for public and private organizations.

Ekrem Tatoglu is Professor of International Business at Gulf University for Science and Technology, Kuwait, and Ibn Haldun University, Turkey. His research focuses on international business, innovation, and operations management. He has published over 140 journal articles and has been a Turkish Academy of Sciences (TÜBA) Fellow since 2015.

Aysun Sahin is Lecturer in the Department of Electronics and Automation at Düzce University, Turkey. Her interdisciplinary research focuses on consumer behavior, digital marketing and generative artificial intelligence. She has academic and industry experience in the United Kingdom and Turkey across engineering and software development.