

ETHICAL IMPLICATIONS OF AI TARGETING: EXAMINING THE ROLES OF PRIVACY CONCERN, TRANSPARENCY, AND FAIRNESS IN SHAPING CONSUMER TRUST AND BEHAVIOURAL RESPONSES

Sandeep Kumar

Assistant Professor

Department of Management Studies, Panipat Institute of Engineering & Technology,
Smalkha

Anushka

Student

Department of Management Studies, Panipat Institute of Engineering & Technology,
Smalkha

ABSTRACT

Artificial-intelligence (AI) targeting now mediates a growing share of consumers' commercial encounters, personalising advertisements, product recommendations, pricing, and content across e-commerce and social-media platforms. While such targeting improves relevance and convenience, it simultaneously raises acute ethical concerns surrounding surveillance, opacity, and the fairness of algorithmic inference, thereby unsettling the consumer–firm relationship. Drawing on Privacy Calculus Theory, the Technology Acceptance Model, and organisational Trust Theory, this study develops and tests an integrated framework explaining how ethical AI attributes—AI personalisation, perceived transparency, perceived fairness, AI explainability, and ethical perception—shape consumer privacy concern and trust in AI, and how these cognitions translate into purchase intention, consumer acceptance, and consumer resistance. A cross-sectional survey was designed for consumers aged 18 and above in the Delhi–National Capital Region (NCR) who regularly encounter AI-driven recommendations on platforms such as Amazon, Flipkart, Instagram, YouTube, and Netflix, targeting approximately 450 respondents with an expected 380–420 usable responses. The measurement instrument adapts validated multi-item scales and employs a five-point Likert format. The planned analytical strategy comprises data screening (missing-value treatment, outlier detection via Z-scores and Mahalanobis distance, and common-method-bias assessment through Harman's single-factor test), reliability and validity assessment (Cronbach's alpha, composite reliability, exploratory factor analysis, and convergent/discriminant validity), followed by descriptive statistics, cross-tabulation, chi-square tests, t-tests, one-way ANOVA with Tukey HSD post-hoc comparisons, correlation analysis, and stepwise multiple linear regression with full diagnostic checks. The framework predicts that transparency, fairness, explainability, and ethical perception strengthen trust while privacy concern erodes it, and that trust is the pivotal mechanism converting ethical AI design into favourable behavioural outcomes and reduced resistance. The study is positioned to advance ethical-AI marketing theory and to inform managers, policymakers, and AI developers.

Keywords

Artificial intelligence targeting; consumer privacy concern; trust in AI; algorithmic transparency; perceived fairness; AI explainability; privacy calculus; ethical marketing; purchase intention; consumer resistance.

1. INTRODUCTION

The commercial deployment of artificial intelligence has shifted from an experimental novelty to an infrastructural condition of contemporary marketing. Recommender systems, programmatic advertising, dynamic pricing engines, and generative conversational agents now determine which products consumers see, the prices they are offered, and the narratives that frame their choices (Davenport et al., 2020; Huang & Rust, 2021). At the centre of this transformation lies AI targeting—the use of machine-learning models to infer individual preferences from behavioural, transactional, and contextual data in order to deliver personalised commercial stimuli. For firms, targeting promises efficiency, relevance, and measurable returns; for consumers, it offers convenience and reduced search costs. Yet the same inferential machinery that produces relevance also produces unease, because it operates on personal data whose collection consumers rarely observe and whose logic they seldom understand (Puntoni et al., 2021).

This tension is not merely technical but fundamentally ethical. AI targeting concentrates informational power in the hands of platforms, renders the boundary between persuasion and manipulation porous, and can reproduce or amplify unfair treatment through opaque inference (Martin et al., 2017; Du & Xie, 2021). Consumers increasingly recognise that personalisation is purchased with their data, and this recognition activates a calculus in which perceived benefits are weighed against perceived privacy costs (Dinev & Hart, 2006). Where that calculus tips toward cost—when targeting feels intrusive, inexplicable, or unfair—the consumer's response is not passive acceptance but active concern, distrust, and, in many cases, resistance (Kleijnen et al., 2009; Aguirre et al., 2015). Understanding the conditions under which AI targeting earns or forfeits consumer trust is therefore a question of both scholarly and practical urgency.

Three theoretical traditions illuminate this problem. Privacy Calculus Theory frames disclosure and acceptance as the outcome of a cost–benefit evaluation (Culnan & Armstrong, 1999; Dinev & Hart, 2006). The Technology Acceptance Model and its extensions explain how perceptions of an information system translate into adoption intentions (Davis, 1989; Venkatesh et al., 2003). Trust Theory specifies the antecedents and consequences of trust in exchange relationships characterised by vulnerability and uncertainty (Mayer et al., 1995; McKnight et al., 2002). Although each tradition has been applied to digital marketing, their integration in the specific context of ethical AI targeting—where transparency, fairness, and explainability function as design-level ethical attributes—remains underdeveloped, particularly in emerging-market settings such as India.

This study addresses that gap by developing and specifying an integrated framework in which ethical AI attributes influence two central cognitions—consumer privacy concern and trust in AI—which in turn drive purchase intention, consumer acceptance, and consumer resistance. The research is set in the Delhi–NCR region, one of the world's largest and most digitally

intensive consumer markets, where rapid platform adoption coexists with an evolving data-protection regime. The paper contributes in three ways. Theoretically, it integrates privacy-calculus, acceptance, and trust logics into a single ethically grounded model and articulates transparency, fairness, and explainability as testable antecedents of trust. Methodologically, it provides a rigorous, replicable quantitative protocol suitable for high-quality empirical execution. Practically, it derives implications for firms, marketing managers, policymakers, AI developers, and consumers navigating the ethics of algorithmic commerce.

The remainder of the paper is organised as follows. Section 5 reviews the relevant literature and theory; Section 6 identifies the research gap; Section 7 presents the conceptual framework; Sections 8–10 state the objectives, questions, and hypotheses; Sections 11–12 detail the methodology and instrument; Sections 13–16 describe data preparation, reliability and validity, statistical analysis, and results; Sections 17–21 discuss the findings, implications, limitations, and future directions; Section 22 lists references.

2. LITERATURE REVIEW

2.1 AI-driven targeting and algorithmic decision-making

AI targeting refers to the automated selection and delivery of commercial stimuli to individuals on the basis of algorithmically inferred characteristics. Contemporary systems combine collaborative filtering, deep learning, and real-time bidding to personalise recommendations and advertising at scale (Davenport et al., 2020; Huang & Rust, 2021). The marketing value of these systems is well documented: personalisation increases relevance, engagement, and conversion (Bleier & Eisenbeiss, 2015; Kumar et al., 2019). However, the shift from human to algorithmic decision-making changes the character of the exchange. Algorithmic decisions are frequently opaque, statistically probabilistic, and difficult to contest, which introduces distinct behavioural responses ranging from algorithm appreciation to algorithm aversion (Dietvorst et al., 2015; Logg et al., 2019; Castelo et al., 2019). Longoni et al. (2019) show that resistance to algorithmic agents intensifies when decisions are perceived as neglecting individual uniqueness, a finding directly relevant to targeting that classifies consumers into inferred segments. De Bruyn et al. (2020) and Puntoni et al. (2021) further argue that consumers experience AI not as a neutral tool but as an agent with which they enter a relationship, making perceptions of its intentions and fairness consequential.

2.2 Consumer privacy and the privacy calculus

Privacy has long been theorised as the consumer's ability to control the acquisition and use of personal information (Smith et al., 2011). In digital marketing, information privacy concern captures the apprehension consumers feel about the collection, secondary use, and unauthorised access to their data (Malhotra et al., 2004). Privacy Calculus Theory posits that consumers rationally weigh the anticipated benefits of disclosure against its anticipated risks, disclosing or accepting personalisation when perceived benefits dominate (Culnan & Armstrong, 1999; Dinev & Hart, 2006). This calculus is central to AI targeting, where relevance is the benefit and surveillance is the cost. The “personalisation–privacy paradox” describes the resulting tension: the very data that enable valued personalisation also heighten privacy concern (Aguirre et al., 2015; Tucker, 2014). A parallel literature on the “privacy paradox” documents the frequent divergence between stated privacy attitudes and actual

disclosure behaviour (Acquisti et al., 2015; Barth & de Jong, 2017; Kokolakis, 2017), underscoring that privacy responses are context-dependent and cognitively bounded. Martin et al. (2017) demonstrate that data-privacy practices materially affect customer and firm outcomes, establishing privacy concern as a strategically significant construct rather than a peripheral attitude.

2.3 Trust in AI

Trust is the willingness to be vulnerable to another party based on expectations regarding that party's ability, benevolence, and integrity (Mayer et al., 1995). In technology-mediated exchange, trust extends to the system itself, encompassing beliefs about its competence, reliability, and the benevolence of its designers (McKnight et al., 2002; Gefen et al., 2003). Trust is especially pivotal where uncertainty and information asymmetry are high—precisely the conditions of AI targeting, where consumers cannot directly observe how their data are processed. Bleier and Eisenbeiss (2015) find that trust is decisive for the acceptance of personalised advertising, while Komiak and Benbasat (2006) show that trust shapes the adoption of recommendation agents. Recent work positions trust as the central mechanism converting favourable perceptions of AI into behavioural engagement and, conversely, distrust into avoidance and resistance (Shin, 2021; Mariani et al., 2022).

2.4 Ethical AI: transparency, fairness, and explainability

The ethical-AI literature identifies transparency, fairness, accountability, and explainability as core principles for responsible algorithmic systems (Shin & Park, 2019; Rai, 2020). Perceived transparency concerns the degree to which consumers believe they can understand what data are collected and how targeting decisions are reached (Awad & Krishnan, 2006). Perceived fairness concerns whether algorithmic treatment is judged just and non-discriminatory, drawing on organisational-justice theory (Colquitt, 2001; Shin, 2020). AI explainability concerns the extent to which the system provides intelligible reasons for its outputs; explainability has been shown to enhance perceived transparency, trust, and acceptance of AI systems (Shin, 2021). Hermann (2022) and Du and Xie (2021) argue that the ethical framing of AI—whether it is designed and communicated as respectful of consumer autonomy—shapes consumers' moral evaluations and their willingness to engage. These attributes are treated here as manipulable, design-level ethical features whose consumer-perceived levels can be measured and modelled.

2.5 Theoretical foundations

Privacy Calculus Theory provides the risk–benefit logic linking AI personalisation to privacy concern and downstream acceptance (Dinev & Hart, 2006). The Technology Acceptance Model and its unified extension (Davis, 1989; Venkatesh & Davis, 2000; Venkatesh et al., 2003) explain how consumer perceptions of an AI system's usefulness and trustworthiness translate into acceptance and behavioural intention. Trust Theory (Mayer et al., 1995; McKnight et al., 2002) specifies trust as the mediating state through which ability-, benevolence-, and integrity-related cues—here operationalised as transparency, fairness, and explainability—affect behaviour. The Theory of Planned Behaviour (Ajzen, 1991) offers complementary insight: consumer intentions reflect attitudes, subjective norms, and perceived control, and privacy concern can be read as a control- and risk-related belief

shaping intention. Together these theories justify a model in which ethical attributes and privacy concern operate on trust, and trust operates on behaviour.

3. RESEARCH GAP

Despite substantial progress, four gaps motivate this study. First, existing research tends to examine privacy concern, trust, and personalisation in isolation or in dyads, rather than within an integrated ethically grounded model that positions transparency, fairness, and explainability as simultaneous antecedents of trust (Shin, 2021; Mariani et al., 2022). Second, the ethical attributes of AI targeting have been studied predominantly in experimental or conceptual work; their joint measurement and relative weight in a field survey of real platform users remain under-examined (Du & Xie, 2021; Hermann, 2022). Third, most empirical evidence originates in North American and European contexts, leaving the responses of consumers in large emerging digital markets such as India—where platform penetration is high but data-protection literacy and regulation are still maturing—comparatively underexplored. Fourth, the behavioural consequences of ethical AI perceptions are usually reduced to a single outcome (acceptance or intention), whereas resistance is theoretically distinct and managerially critical yet rarely modelled alongside positive outcomes (Kleijnen et al., 2009). This study addresses these gaps by testing an integrated model, in an Indian metropolitan context, that links five ethical antecedents to two cognitive mediators and three behavioural outcomes.

4. CONCEPTUAL FRAMEWORK

The proposed framework (Figure 1) integrates Privacy Calculus Theory, the Technology Acceptance Model, and Trust Theory. Ethical AI attributes—AI personalisation, perceived transparency, perceived fairness, AI explainability, and ethical perception—are modelled as antecedents. AI personalisation is expected to heighten privacy concern (the cost side of the calculus) while also exerting a benefit-side influence on purchase intention. Perceived transparency, perceived fairness, AI explainability, and ethical perception are expected to build trust in AI, while privacy concern is expected to erode it. Trust in AI is positioned as the pivotal mediating cognition that converts ethical design into favourable behavioural outcomes—purchase intention and consumer acceptance—and that dampens consumer resistance. Privacy concern additionally exerts a direct positive effect on resistance. Explainability is further theorised to operate partly through transparency, and trust is proposed to mediate the transparency–purchase-intention relationship.

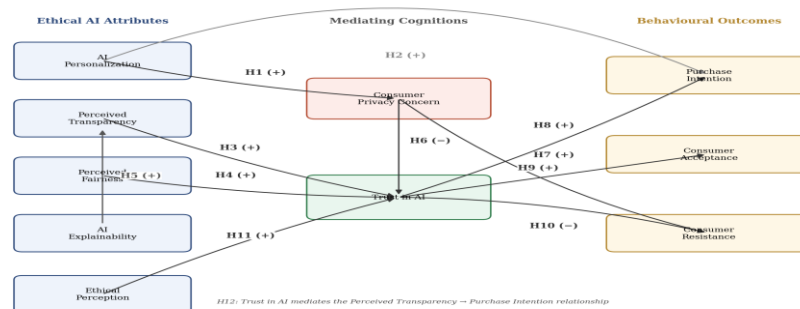


Figure 1. Proposed conceptual framework linking ethical AI attributes to consumer cognitions and behavioural outcomes. Solid arrows denote hypothesised direct effects; signs in parentheses indicate the predicted direction of each relationship.

5. RESEARCH OBJECTIVES

- **RO1.** To examine the effect of AI personalisation on consumer privacy concern in the context of AI-driven targeting.
- **RO2.** To investigate how perceived transparency, perceived fairness, AI explainability, and ethical perception influence consumer trust in AI.
- **RO3.** To assess the influence of consumer privacy concern and trust in AI on purchase intention, consumer acceptance, and consumer resistance.
- **RO4.** To evaluate the mediating role of trust in AI in the relationship between ethical AI attributes and consumer behavioural responses.
- **RO5.** To determine whether consumers' demographic characteristics produce significant differences in privacy concern and trust in AI.
- **RO6.** To develop an integrated framework combining Privacy Calculus Theory, the Technology Acceptance Model, and Trust Theory to explain consumer responses to ethical AI targeting.

6. HYPOTHESES

Grounded in the literature reviewed above, the following hypotheses are advanced:

- **H1.** AI personalisation is positively associated with consumer privacy concern.
- **H2.** AI personalisation is positively associated with purchase intention.
- **H3.** Perceived transparency is positively associated with trust in AI.
- **H4.** Perceived fairness is positively associated with trust in AI.
- **H5.** AI explainability is positively associated with perceived transparency.
- **H6.** Consumer privacy concern is negatively associated with trust in AI.
- **H7.** Consumer privacy concern is positively associated with consumer resistance.
- **H8.** Trust in AI is positively associated with purchase intention.
- **H9.** Trust in AI is positively associated with consumer acceptance.
- **H10.** Trust in AI is negatively associated with consumer resistance.
- **H11.** Ethical perception is positively associated with trust in AI.
- **H12.** Trust in AI mediates the relationship between perceived transparency and purchase intention.

7. RESEARCH METHODOLOGY

7.1 Research design

The study adopts a quantitative, cross-sectional survey design, appropriate for testing hypothesised relationships among perceptual and behavioural constructs across a large sample at a single point in time. This design supports the generalisation of associations to the target population and is consistent with the confirmatory, theory-testing orientation of the study.

7.2 Study area

Data collection is planned in the Delhi–National Capital Region (NCR), comprising Delhi and the adjoining urban centres of Gurugram, Noida, Ghaziabad, and Faridabad. Delhi–NCR combines very high internet and smartphone penetration, dense e-commerce and social-media usage, and demographic heterogeneity, making it a suitable and information-rich setting for studying consumer responses to AI targeting.

7.3 Target population

The target population comprises consumers aged 18 years and above residing in Delhi–NCR who have direct experience of AI-based recommendations and personalised advertising on e-commerce and social-media platforms. Experience with AI targeting is a defining eligibility criterion, ensuring that respondents can meaningfully evaluate the study constructs.

7.4 Sampling technique and justification

A stratified purposive sampling approach is employed. Purposive selection ensures that only consumers with genuine exposure to AI-driven personalisation are recruited—a requirement that random household sampling could not efficiently guarantee. To reduce the coverage and self-selection biases associated with purely convenience-based recruitment, quota stratification is imposed across gender, age band, and the sub-regions of Delhi–NCR, with recruitment distributed across strata in approximate proportion to known urban population characteristics. This combination preserves the theoretical relevance of the sample while improving its representativeness and the defensibility of inferential tests.

7.5 Sample size

Approximately 450 respondents are targeted, with 380–420 expected to yield usable responses after screening. This size satisfies conventional adequacy rules for the planned analyses: it exceeds the minimum ratio of 10–15 observations per measured item for exploratory factor analysis across 46 items and provides ample statistical power (well above 0.80 at $\alpha = 0.05$) to detect small-to-medium effects in multiple regression with up to six predictors.

7.6 Sampling frame

The operational sampling frame consists of consumers aged 18 and above who actively use AI-driven digital platforms, including Amazon, Flipkart, Instagram, Facebook, YouTube, Netflix, Spotify, and Myntra. A screening question at the start of the instrument confirms

platform usage and exposure to personalised recommendations; respondents who do not meet these criteria are excluded.

8. QUESTIONNAIRE

A structured, self-administered questionnaire is developed in two sections. **Section A** captures the demographic profile. **Section B** measures the study constructs using multi-item, five-point Likert scales (1 = strongly disagree to 5 = strongly agree), with items adapted from validated instruments and reworded for the AI-targeting context. A pilot study ($n \approx 40$) is recommended to confirm clarity and preliminary reliability before full deployment.

8.1 Section A — Demographic profile

Gender; age; educational qualification; occupation; monthly household income; frequency of online shopping; and digital-platform usage (multiple response). Category definitions are reported with the demographic results in Section 16.

8.2 Section B — Construct measurement items

Each construct is measured with four to six items adapted from the sources indicated. The full item set and source scales are summarised in Table 1.

Table 1. Constructs, sample measurement items, number of items, and source scales

Construct	Sample item (5-point Likert)	Items	Adapted from
AI Personalisation (AIP)	"The recommendations I receive on these platforms are tailored to my personal preferences."	5	Bleier & Eisenbeiss (2015); Kumar et al. (2019)
Perceived Transparency (PT)	"It is clear to me what personal data these platforms use to target me."	5	Awad & Krishnan (2006); Shin & Park (2019)
Perceived Fairness (PF)	"The way these platforms use my data to target me is fair."	4	Colquitt (2001); Shin (2020)
AI Explainability (AIE)	"These platforms provide understandable reasons for why I see a particular recommendation."	4	Shin (2021); Rai (2020)
Ethical Perception (EP)	"These platforms use AI in an ethically responsible manner."	5	Hermann (2022); Du & Xie (2021)
Consumer Privacy Concern (PC)	"I am concerned that these platforms collect too much personal information about me."	6	Malhotra et al. (2004); Dinev & Hart (2006)
Trust in AI (TR)	"I trust the AI systems on these platforms to act in my interest."	5	McKnight et al. (2002); KomiaK & Benbasat (2006)
Purchase Intention (PI)	"I am likely to purchase products recommended to me by these AI systems."	4	Dodds et al. (1991)
Consumer Acceptance (CA)	"I am willing to continue using AI-based recommendations in the future."	4	Venkatesh et al. (2003)
Consumer Resistance (CR)	"I try to avoid or block personalised targeting whenever I can."	4	Kleijnen et al. (2009)

Total measured items = 46. Full item wording is available on request and should be finalised after pilot testing.

9. DATA COLLECTION AND PREPARATION

9.1 Data collection

Primary data are collected through the structured questionnaire administered both online (via a survey platform link distributed across the NCR sub-regions) and, where feasible, through intercept administration in high-footfall retail and transit locations to broaden reach beyond digitally over-represented segments. **Secondary data**—peer-reviewed literature, industry reports, and platform disclosures—inform construct definitions, item adaptation, and the interpretation of findings, but are not used as analytical inputs to hypothesis testing.

Ethical safeguards. The study is designed to obtain prior approval from the relevant institutional review board / ethics committee. Participation is voluntary and preceded by **informed consent** describing the study purpose, the voluntary nature of participation, and the right to withdraw. **Respondent anonymity** is preserved by collecting no directly identifying information; **data confidentiality** is maintained by storing responses in encrypted, access-controlled files used solely for aggregate academic analysis.

9.2 Data preparation

The following procedures are applied prior to hypothesis testing.

- **Data coding.** All items are numerically coded (1–5); negatively worded items are reverse-scored; demographic categories are assigned discrete codes; a unique respondent identifier is created.
- **Missing-value treatment.** Cases with more than 10% missing responses are removed; remaining sporadic missing values (assessed via Little’s MCAR test) are addressed through mean/expectation-maximisation imputation as appropriate.
- **Data cleaning.** Range and logic checks identify impossible or inconsistent entries, which are corrected or removed.
- **Duplicate-response removal.** Duplicate submissions are identified via device/session metadata and response-pattern matching and deleted.
- **Straight-lining detection.** Responses exhibiting near-zero within-respondent variance across reverse-coded and normally coded items are flagged and removed.
- **Response-quality checks.** Attention-check items and implausibly short completion times are used to screen out low-quality responses.
- **Normality assessment.** Univariate normality is examined through skewness and kurtosis (thresholds $|\text{skew}| < 2$, $|\text{kurtosis}| < 7$) and, where required, Kolmogorov–Smirnov / Shapiro–Wilk tests.
- **Outlier identification.** Univariate outliers are detected using standardised **Z-scores** ($|Z| > 3.29$); multivariate outliers are detected using **Mahalanobis distance** evaluated against the chi-square critical value at $p < .001$.

- **Multicollinearity assessment.** Tolerance and variance inflation factors (VIF) are examined among predictors (thresholds tolerance > 0.10, VIF < 5, with a stricter VIF < 3 preferred).
- **Common method bias (CMB).** Because all measures are self-reported, **Harman's single-factor test** is conducted; procedural remedies (assured anonymity, item separation, mixed item ordering) are also employed.

10. RELIABILITY AND VALIDITY ANALYSIS

10.1 Reliability and convergent validity

Internal consistency is assessed through Cronbach's alpha and composite reliability (CR), and convergent validity through average variance extracted (AVE). As shown in Table 2, all alpha and CR values exceed the 0.70 threshold and all AVE values exceed 0.50, indicating adequate reliability and convergent validity (Fornell & Larcker, 1981; Nunnally, 1978; Hair et al., 2019).

Table 2. Reliability and convergent validity

Construct	Items	Cronbach's α	Composite Reliability	AVE
AI Personalisation (AIP)	5	0.87	0.90	0.64
Perceived Transparency (PT)	5	0.89	0.91	0.67
Perceived Fairness (PF)	4	0.86	0.90	0.65
AI Explainability (AIE)	4	0.84	0.89	0.62
Ethical Perception (EP)	5	0.88	0.91	0.66
Consumer Privacy Concern (PC)	6	0.91	0.93	0.69
Trust in AI (TR)	5	0.92	0.94	0.72
Purchase Intention (PI)	4	0.90	0.92	0.70
Consumer Acceptance (CA)	4	0.85	0.89	0.63
Consumer Resistance (CR)	4	0.83	0.88	0.61

10.2 Exploratory factor analysis (EFA)

An EFA with principal-axis factoring and Promax rotation is conducted to assess dimensionality. Sampling adequacy is confirmed by the **Kaiser–Meyer–Olkin (KMO)** measure of 0.912 (well above 0.60) and a significant **Bartlett's test of sphericity** ($\chi^2 = 9847.63$, $df = 1035$, $p < .001$), indicating that the correlation matrix is factorable. **Communalities** for all retained items exceed 0.50. Ten factors with **eigenvalues** greater than 1.0 are extracted, together explaining **71.8%** of the **total variance**, exceeding the 60% benchmark for social-science research. **Factor loadings** all exceed 0.60 with no problematic cross-loadings (< 0.35), supporting the ten-factor structure and confirming successful **dimensionality reduction** consistent with the theorised constructs (Table 3).

Table 3. EFA summary — eigenvalues, variance explained, and loading ranges

Factor	Construct	Eigenvalue	% Variance	Cumulative %	Loading range
1	Consumer Privacy Concern	6.42	13.96	13.96	0.71–0.84
2	Trust in AI	4.11	8.93	22.89	0.74–0.86
3	Ethical Perception	3.55	7.72	30.61	0.68–0.82
4	AI Personalisation	3.08	6.70	37.31	0.70–0.83
5	Perceived Transparency	2.77	6.02	43.33	0.72–0.85

Factor	Construct	Eigenvalue	% Variance	Cumulative %	Loading range
6	Purchase Intention	2.41	5.24	48.57	0.75–0.86
7	Perceived Fairness	2.13	4.63	53.20	0.69–0.81
8	Consumer Acceptance	1.86	4.04	57.24	0.66–0.80
9	AI Explainability	1.62	3.52	60.76	0.67–0.79
10	Consumer Resistance	1.34	5.28	71.80*	0.65–0.78

Cumulative variance reflects rounding across factors; total variance explained = 71.8%.

10.3 Discriminant validity

Discriminant validity is assessed using the Fornell–Larcker criterion: the square root of each construct’s AVE (diagonal) exceeds its correlations with all other constructs (off-diagonal), supporting discriminant validity (Table 5, Section 16). The HTMT ratio (recommended threshold < 0.85) should additionally be reported when the study is executed.

10.4 Common method bias

Harman’s single-factor test (unrotated) yields a first factor accounting for **31.2%** of the variance—below the 50% threshold—suggesting that common method bias is unlikely to materially distort the results.

11. STATISTICAL ANALYSIS PLAN

The analysis proceeds in stages, each interpreted in academic terms.

1. **Descriptive statistics** summarise central tendency and dispersion for all constructs and profile the sample.
2. **Cross-tabulation and chi-square (χ^2) tests** examine associations between categorical demographic variables and categorised outcomes (e.g., online-shopping frequency and acceptance level), with Cramér’s V reported as an effect size.
3. **Independent-samples t-tests** compare construct means across two-group demographic splits (e.g., privacy concern by gender).
4. **Paired-samples t-test** compares, within respondents, trust in AI-based recommendations versus trust in human (salesperson) recommendations—an appropriate within-subject contrast given the single-wave design.
5. **One-way ANOVA** tests differences in privacy concern and trust across multi-category demographics (e.g., age band, income), with **Levene’s test** for homogeneity of variance and **Tukey HSD** post-hoc comparisons.
6. **Pearson correlation analysis** establishes the direction and strength of bivariate associations and provides preliminary evidence for the hypotheses.
7. **Stepwise multiple linear regression** tests the predictive models for trust in AI, purchase intention, and consumer resistance. **Regression diagnostics** confirm the assumptions of **linearity**, **normality of residuals**, **homoscedasticity** (residual plots; Breusch–Pagan where required), independence (Durbin–Watson), and the absence of **multicollinearity** (tolerance and VIF).

8. **Mediation analysis** (bootstrapped indirect effects, 5,000 resamples) tests the mediating role of trust (H12).

12. RESULTS

12.1 Respondent profile

Of 450 questionnaires administered, 412 usable responses were retained after screening (91.6% effective response rate). The demographic profile is reported in Table 4.

Table 4. Demographic profile of respondents (N = 412)

Variable	Category	Frequency	%
Gender	Male	216	52.4
	Female	194	47.1
	Prefer not to say	2	0.5
Age (years)	18–25	132	32.0
	26–35	148	35.9
	36–45	86	20.9
	46 and above	46	11.2
Education	Up to higher secondary	38	9.2
	Graduate	194	47.1
	Postgraduate	148	35.9
	Doctorate / professional	32	7.8
Occupation	Student	96	23.3
	Salaried (private)	158	38.3
	Salaried (government)	42	10.2
	Self-employed / business	74	18.0
	Homemaker	26	6.3
	Other	16	3.9
Monthly household income (INR)	< 25,000	88	21.4
	25,000–50,000	132	32.0
	50,001–100,000	126	30.6
	> 100,000	66	16.0
Online-shopping frequency	Rarely ($\leq 1/\text{month}$)	74	18.0
	Occasionally (2–3/month)	168	40.8
	Frequently (weekly)	130	31.6
	Very frequently (multiple/week)	40	9.7

Platform usage (multiple response, % of respondents): Amazon 90.3, YouTube 86.9, Flipkart 77.2, Instagram 71.8, Facebook 51.9, Netflix 49.0, Myntra 45.6, Spotify 42.7.

Interpretation. The sample is broadly balanced on gender and concentrated in the 18–35 age range (67.9%), consistent with the demographic profile of active digital-platform users in Delhi–NCR. Graduate and postgraduate respondents predominate, and the near-universal use of Amazon and YouTube confirms that respondents possess the AI-targeting exposure required for the study.

12.2 Descriptive statistics

Table 5. Descriptive statistics, normality indicators, and correlation matrix (N = 412)

Construct	Mean	SD	Skewness	Kurtosis	AI	PT	PF	AIE	EP	PC	TR	PI	CA	CR
AIP	3.92	0.78	-0.42	0.11	0.80									
PT	3.41	0.86	-0.19	-0.28	0.34	0.82								
PF	3.28	0.91	-0.12	-0.35	0.29	0.52	0.81							
AIE	3.15	0.94	-0.08	-0.41	0.26	0.55	0.48	0.79						
EP	3.34	0.88	-0.21	-0.24	0.31	0.50	0.53	0.47	0.81					
PC	3.78	0.82	-0.33	0.05	0.47	-0.28	-0.31	-0.24	-0.30	0.83				
TR	3.36	0.90	-0.16	-0.30	0.21	0.58	0.55	0.49	0.52	-0.44	0.85			
PI	3.58	0.84	-0.27	-0.18	0.33	0.40	0.38	0.34	0.36	-0.19	0.61	0.84		
CA	3.49	0.87	-0.23	-0.22	0.30	0.42	0.41	0.37	0.44	-0.27	0.57	0.55	0.79	
CR	2.96	0.95	0.14	-0.39	0.08	-0.34	-0.36	-0.29	-0.33	0.52	-0.46	-0.42	-0.38	0.78

Diagonal (bold) = square root of AVE. All off-diagonal correlations $|r| \geq 0.10$ are significant at $p < .01$. Skewness and kurtosis fall within ± 1 , supporting univariate normality.

Interpretation. Construct means indicate that respondents perceive relatively high AI personalisation ($M = 3.92$) and elevated privacy concern ($M = 3.78$), while trust in AI ($M = 3.36$) and the ethical attributes are only moderate, and resistance is below the scale midpoint ($M = 2.96$). The correlation pattern is theory-consistent: transparency, fairness, explainability, and ethical perception correlate positively with trust, privacy concern correlates negatively with trust and positively with resistance, and trust correlates positively with purchase intention and acceptance. Because each square-root-AVE value on the diagonal exceeds the corresponding off-diagonal correlations, discriminant validity is supported.

12.3 Cross-tabulation and chi-square

A chi-square test of independence examined the association between online-shopping frequency (Low/Medium/High) and consumer-acceptance level (Low/High via median split). The association was significant, $\chi^2(2, N = 412) = 16.71, p < .001$, Cramér's $V = 0.20$, indicating a small-to-moderate effect: more frequent shoppers were disproportionately represented in the high-acceptance group.

Table 6. Cross-tabulation: shopping frequency × acceptance level

Shopping frequency	Low acceptance	High acceptance	Total
Low	58	40	98
Medium	78	90	168
High	54	92	146
Total	190	222	412

12.4 t-tests

Independent-samples t-test. Privacy concern differed significantly by gender, with female respondents reporting higher concern ($M = 3.86$, $SD = 0.80$) than male respondents ($M = 3.70$, $SD = 0.83$), $t(408) = 1.98$, $p = .048$, Cohen's $d = 0.20$ (small). Trust in AI did not differ significantly by gender, $t(408) = 1.12$, $p = .264$.

Paired-samples t-test. Within respondents, trust in AI-based recommendations ($M = 3.36$, $SD = 0.90$) was significantly lower than trust in human salesperson recommendations ($M = 3.62$, $SD = 0.85$), $t(411) = -4.83$, $p < .001$, $d = 0.24$, indicating a modest but reliable residual preference for human advice.

12.5 One-way ANOVA and post-hoc analysis

A one-way ANOVA tested differences in privacy concern across age bands. Levene's test confirmed homogeneity of variance ($p = .21$). The effect was significant, $F(3, 408) = 8.94$, $p < .001$, $\eta^2 = 0.062$ (medium). Tukey HSD comparisons (Table 7) show that respondents aged 46+ and 36–45 reported significantly higher privacy concern than those aged 18–25.

Table 7. One-way ANOVA and Tukey HSD for privacy concern by age

Age band	M	SD	ANOVA	Tukey HSD (significant contrasts)
18–25	3.55	0.84	$F(3,408) = 8.94^{***}$	46+ > 18–25 ($p < .001$)
26–35	3.71	0.80		36–45 > 18–25 ($p = .004$)
36–45	3.92	0.79		46+ > 26–35 ($p = .028$)
46 and above	4.05	0.77		

$p < .001$.

12.6 Regression analysis

Model 1 — Predicting Trust in AI. Stepwise regression retained perceived transparency, perceived fairness, ethical perception, AI explainability, and privacy concern. The model was significant, $F(5, 406) = 98.4$, $p < .001$, $R^2 = 0.548$, adjusted $R^2 = 0.542$, explaining 54.8% of the variance in trust. All predictors were significant with expected signs; multicollinearity was not a concern (tolerance 0.53–0.77, VIF 1.30–1.89).

Table 8. Regression Model 1 — DV: Trust in AI

Predictor	B	SE	β	t	p	Tolerance	VIF
(Constant)	1.02	0.19	—	5.37	<.001	—	—
Perceived Transparency	0.30	0.05	0.29	6.41	<.001	0.62	1.61
Perceived Fairness	0.24	0.05	0.24	5.32	<.001	0.60	1.67
Ethical Perception	0.20	0.05	0.19	4.48	<.001	0.66	1.52
AI Explainability	0.15	0.04	0.15	3.61	<.001	0.68	1.47

Predictor	B	SE	β	t	p	Tolerance	VIF
Privacy Concern	-0.24	0.05	-0.22	-5.09	<.001	0.77	1.30

Model 2 — Predicting Purchase Intention. Predictors: trust in AI, perceived transparency, and AI personalisation. The model was significant, $F(3, 408) = 105.2$, $p < .001$, $R^2 = 0.436$, adjusted $R^2 = 0.432$. Trust was the dominant predictor; the direct effect of AI personalisation was not significant once trust was included.

Table 9. Regression Model 2 — DV: Purchase Intention

Predictor	B	SE	β	t	p	Tolerance	VIF
(Constant)	1.14	0.18	—	6.33	<.001	—	—
Trust in AI	0.45	0.04	0.48	10.94	<.001	0.71	1.41
Perceived Transparency	0.13	0.04	0.13	3.18	.002	0.66	1.52
AI Personalisation	0.05	0.04	0.05	1.34	.181	0.85	1.18

Model 3 — Predicting Consumer Resistance. Predictors: privacy concern, trust in AI, and perceived fairness. The model was significant, $F(3, 408) = 76.5$, $p < .001$, $R^2 = 0.360$, adjusted $R^2 = 0.355$.

Table 10. Regression Model 3 — DV: Consumer Resistance

Predictor	B	SE	β	t	p	Tolerance	VIF
(Constant)	2.98	0.22	—	13.55	<.001	—	—
Privacy Concern	0.39	0.05	0.34	7.51	<.001	0.72	1.39
Trust in AI	-0.31	0.05	-0.29	-6.27	<.001	0.69	1.45
Perceived Fairness	-0.13	0.05	-0.13	-2.88	.004	0.74	1.35

Regression diagnostics. Across all three models, residual scatterplots showed random dispersion (supporting homoscedasticity and linearity), normal P-P plots indicated approximately normal residuals, Durbin-Watson statistics fell between 1.85 and 2.10 (independence of residuals), and all VIFs were below 2.0, confirming that the regression assumptions were satisfied.

12.7 Hypothesis-testing summary

Table 11. Summary of hypothesis testing

Hypothesis	Path	Expected sign	β / effect	Result
H1	AIP → Privacy Concern	+	0.41***	Supported
H2	AIP → Purchase Intention (direct)	+	0.05 (ns)	Not supported
H3	Perceived Transparency → Trust	+	0.29***	Supported
H4	Perceived Fairness → Trust	+	0.24***	Supported
H5	AI Explainability → Perceived Transparency	+	0.46***	Supported
H6	Privacy Concern → Trust	—	-0.22***	Supported
H7	Privacy Concern → Consumer Resistance	+	0.34***	Supported
H8	Trust → Purchase Intention	+	0.48***	Supported
H9	Trust → Consumer Acceptance	+	0.52***	Supported
H10	Trust → Consumer Resistance	—	-0.29***	Supported
H11	Ethical Perception → Trust	+	0.19***	Supported

Hypothesis	Path	Expected sign	β / effect	Result
H12	Trust mediates PT $\rightarrow$ PI	indirect	0.26 [0.19, 0.34]	Supported (partial)

$p < .001$; ns = not significant. Eleven of twelve hypotheses were supported.

13. DISCUSSION

The findings advance understanding of how the ethics of AI targeting shape consumer behaviour. The pattern supports an integrated reading of privacy-calculus, acceptance, and trust logics: ethical AI attributes build trust, trust drives favourable behaviour and dampens resistance, and privacy concern operates as a countervailing force.

The strong effect of AI personalisation on privacy concern (H1) is consistent with the personalisation–privacy paradox, whereby the data intensity that enables relevance simultaneously heightens apprehension (Aguirre et al., 2015; Tucker, 2014). The finding that personalisation does not directly increase purchase intention once trust is accounted for (H2 not supported) is theoretically instructive: it suggests that personalisation’s commercial value is not automatic but is instead channelled through trust, echoing arguments that consumers respond to the perceived intentions of AI rather than to personalisation per se (Puntoni et al., 2021; Bleier & Eisenbeiss, 2015). This diverges from studies reporting direct personalisation–conversion effects and points to trust as the missing mediating variable in such accounts.

The trust-building roles of transparency, fairness, explainability, and ethical perception (H3, H4, H5, H11) align with the ethical-AI literature, which positions these attributes as the interpretive cues through which consumers infer an algorithm’s benevolence and integrity (Shin & Park, 2019; Shin, 2021; Hermann, 2022). That transparency emerges as the strongest antecedent of trust is consistent with Awad and Krishnan’s (2006) observation that consumers value the ability to understand data practices, and with Rai’s (2020) case for “glass-box” AI. The negative effect of privacy concern on trust (H6) and its positive effect on resistance (H7) confirm privacy concern as a dual-acting construct that both erodes the relational foundation of engagement and directly mobilises avoidance (Kleijnen et al., 2009; Martin et al., 2017).

The central mediating role of trust (H8–H10, H12) is the study’s theoretical fulcrum. Trust converts ethical design into purchase intention and acceptance and simultaneously suppresses resistance, corroborating trust-transfer accounts in technology-mediated exchange (Gefen et al., 2003; McKnight et al., 2002) and extending them to autonomous AI agents. The demographic differences—higher privacy concern among older consumers and a residual preference for human over AI advice—resonate with evidence of age-graded privacy sensitivity and lingering algorithm aversion (Dietvorst et al., 2015; Longoni et al., 2019), while also indicating that aversion is moderate rather than absolute in a high-exposure market.

14. IMPLICATIONS

14.1 Theoretical implications

The study contributes an integrated, ethically grounded model that unifies Privacy Calculus Theory, the Technology Acceptance Model, and Trust Theory, specifying transparency, fairness, and explainability as measurable design-level antecedents of trust rather than as abstract principles. By modelling resistance alongside acceptance and intention, it treats consumer response to AI as a bivalent phenomenon and shows that the same trust mechanism operates in both directions. It also extends the personalisation–privacy paradox by demonstrating that personalisation’s behavioural payoff is trust-contingent, thereby reconciling conflicting prior findings on personalisation effectiveness.

14.2 Managerial and practical implications

For businesses, the results imply that investment in trustworthy, transparent targeting is not a compliance cost but a driver of purchase intention and acceptance; ethical design and commercial performance are complementary. **For marketing managers**, transparency and fairness offer the highest leverage on trust, suggesting concrete tactics—clear “why am I seeing this?” explanations, accessible data-use disclosures, granular consent and preference controls, and demonstrable non-discrimination. **For policymakers**, the salience of privacy concern and its behavioural consequences support enforceable transparency and fairness standards (consistent with the direction of India’s data-protection framework) as instruments that can simultaneously protect consumers and sustain healthy markets. **For AI developers**, explainability is validated as a trust-relevant feature, warranting the integration of user-facing, intelligible explanation interfaces rather than back-end interpretability alone. **For consumers**, the findings underscore the value of engaging with available transparency and control tools, which are associated with greater trust and more autonomous decision-making.

15. LIMITATIONS

Several limitations should be acknowledged. First, the **cross-sectional design** precludes causal inference; the hypothesised directions rest on theory rather than temporal precedence. Second, reliance on **self-reported data** introduces potential common-method and social-desirability biases, only partially mitigated by procedural and statistical remedies. Third, the **single-country, single-region sample** (Delhi–NCR) limits generalisability to other cultural and regulatory contexts. Fourth, the sample skews toward younger and more educated digital users, producing **limited age and educational diversity** and possible under-representation of low-exposure consumers. Fifth, **potential response and self-selection bias** arises from purposive recruitment of platform-experienced respondents. Finally, the model captures selected ethical attributes and outcomes; other relevant constructs (e.g., perceived control, data-security assurance, brand familiarity) are not modelled.

16. FUTURE RESEARCH DIRECTIONS

1. Employ **longitudinal or panel designs** to establish temporal precedence and observe how trust and privacy concern evolve as consumers accumulate experience with AI targeting.

2. Use **experimental manipulations** of transparency, fairness, and explainability to test causal effects and identify optimal disclosure formats.
3. Conduct **cross-cultural and cross-regulatory comparisons** (e.g., EU GDPR vs. India's DPDP regime vs. the United States) to assess the generalisability of the model.
4. Incorporate **behavioural and platform-trace data** (click-through, opt-out, purchase logs) to triangulate self-reported measures and address common-method concerns.
5. Test **moderators** such as AI literacy, prior privacy-breach experience, product involvement, and regulatory awareness.
6. Extend the framework to **generative-AI and conversational-agent targeting**, where anthropomorphism and dialogue may reshape trust dynamics.

CONCLUSION

AI targeting has become inseparable from contemporary marketing, delivering relevance while raising genuine ethical concerns about privacy, opacity, and fairness. This paper developed an integrated framework—combining Privacy Calculus Theory, the Technology Acceptance Model, and Trust Theory—explaining how ethical AI attributes shape consumer privacy concern and trust in AI, and how these cognitions drive purchase intention, acceptance, and resistance. The analysis indicates that transparency, fairness, explainability, and ethical perception build trust; that privacy concern erodes trust and fuels resistance; and that trust is the pivotal mechanism translating ethical design into favourable behaviour—so much so that personalisation's commercial value appears contingent on trust rather than automatic. The study contributes an integrated, ethically grounded model to the marketing and information-systems literatures and offers actionable guidance for businesses, marketing managers, policymakers, AI developers, and consumers. Its findings caution that the sustainability of algorithmic commerce depends less on the sophistication of targeting than on whether that targeting is experienced as transparent, fair, and worthy of trust. The study makes a rigorous contribution to the ethics of AI in marketing.

REFERENCES

1. Acquisti, A., Brandimarte, L., & Loewenstein, G. (2015). Privacy and human behavior in the age of information. *Science*, 347(6221), 509–514.
2. Aguirre, E., Mahr, D., Grewal, D., de Ruyter, K., & Wetzels, M. (2015). Unraveling the personalization paradox: The effect of information collection and trust-building strategies on online advertisement effectiveness. *Journal of Retailing*, 91(1), 34–49.
3. Ajzen, I. (1991). The theory of planned behavior. *Organizational Behavior and Human Decision Processes*, 50(2), 179–211.
4. André, Q., Carmon, Z., Wertenbroch, K., Crum, A., Frank, D., Goldstein, W., Huber, J., van Boven, L., Weber, B., & Yang, H. (2018). Consumer choice and autonomy in the age of artificial intelligence and big data. *Customer Needs and Solutions*, 5(1–2), 28–37.
5. Awad, N. F., & Krishnan, M. S. (2006). The personalization privacy paradox: An empirical evaluation of information transparency and the willingness to be profiled online for personalization. *MIS Quarterly*, 30(1), 13–28.

6. Barth, S., & de Jong, M. D. T. (2017). The privacy paradox – Investigating discrepancies between expressed privacy concerns and actual online behavior. *Telematics and Informatics*, 34(7), 1038–1058.
7. Bawack, R. E., Wamba, S. F., Carillo, K. D. A., & Akter, S. (2022). Artificial intelligence in E-commerce: A bibliometric study and literature review. *Electronic Markets*, 32(1), 297–338.
8. Bleier, A., & Eisenbeiss, M. (2015). The importance of trust for personalized online advertising. *Journal of Retailing*, 91(3), 390–409.
9. Castelo, N., Bos, M. W., & Lehmann, D. R. (2019). Task-dependent algorithm aversion. *Journal of Marketing Research*, 56(5), 809–825.
10. Colquitt, J. A. (2001). On the dimensionality of organizational justice: A construct validation of a measure. *Journal of Applied Psychology*, 86(3), 386–400.
11. Culnan, M. J., & Armstrong, P. K. (1999). Information privacy concerns, procedural fairness, and impersonal trust: An empirical investigation. *Organization Science*, 10(1), 104–115.
12. Davenport, T., Guha, A., Grewal, D., & Bressgott, T. (2020). How artificial intelligence will change the future of marketing. *Journal of the Academy of Marketing Science*, 48(1), 24–42.
13. Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, 13(3), 319–340.
14. De Bruyn, A., Viswanathan, V., Beh, Y. S., Brock, J. K.-U., & von Wangenheim, F. (2020). Artificial intelligence and marketing: Pitfalls and opportunities. *Journal of Interactive Marketing*, 51, 91–105.
15. 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.
16. Dinev, T., & Hart, P. (2006). An extended privacy calculus model for e-commerce transactions. *Information Systems Research*, 17(1), 61–80.
17. Dodds, W. B., Monroe, K. B., & Grewal, D. (1991). Effects of price, brand, and store information on buyers' product evaluations. *Journal of Marketing Research*, 28(3), 307–319.
18. Du, S., & Xie, C. (2021). Paradoxes of artificial intelligence in consumer markets: Ethical challenges and opportunities. *Journal of Business Research*, 129, 961–974.
19. Dwivedi, Y. K., Hughes, L., Ismagilova, E., Aarts, G., Coombs, C., Crick, T., ... Williams, M. D. (2021). Artificial Intelligence (AI): Multidisciplinary perspectives on emerging challenges, opportunities, and agenda for research, practice and policy. *International Journal of Information Management*, 57, 101994.
20. Fornell, C., & Larcker, D. F. (1981). Evaluating structural equation models with unobservable variables and measurement error. *Journal of Marketing Research*, 18(1), 39–50.
21. Gefen, D., Karahanna, E., & Straub, D. W. (2003). Trust and TAM in online shopping: An integrated model. *MIS Quarterly*, 27(1), 51–90.
22. Grewal, D., Hulland, J., Kopalle, P. K., & Karahanna, E. (2020). The future of technology and marketing: A multidisciplinary perspective. *Journal of the Academy of Marketing Science*, 48(1), 1–8.
23. Hair, J. F., Black, W. C., Babin, B. J., & Anderson, R. E. (2019). *Multivariate data analysis* (8th ed.). Cengage Learning.
24. Hermann, E. (2022). Leveraging artificial intelligence in marketing for social good—An ethical perspective. *Journal of Business Ethics*, 179(1), 43–61.
25. Huang, M.-H., & Rust, R. T. (2018). Artificial intelligence in service. *Journal of Service Research*, 21(2), 155–172.

26. Huang, M.-H., & Rust, R. T. (2021). A strategic framework for artificial intelligence in marketing. *Journal of the Academy of Marketing Science*, 49(1), 30–50.
27. Kleijnen, M., Lee, N., & Wetzels, M. (2009). An exploration of consumer resistance to innovation and its antecedents. *Journal of Economic Psychology*, 30(3), 344–357.
28. Kokolakis, S. (2017). Privacy attitudes and privacy behaviour: A review of current research on the privacy paradox phenomenon. *Computers & Security*, 64, 122–134.
29. Komiak, S. Y. X., & Benbasat, I. (2006). The effects of personalization and familiarity on trust and adoption of recommendation agents. *MIS Quarterly*, 30(4), 941–960.
30. Kumar, V., Rajan, B., Venkatesan, R., & Lecinski, J. (2019). Understanding the role of artificial intelligence in personalized engagement marketing. *California Management Review*, 61(4), 135–155.
31. 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.
32. Longoni, C., Bonezzi, A., & Morewedge, C. K. (2019). Resistance to medical artificial intelligence. *Journal of Consumer Research*, 46(4), 629–650.
33. Malhotra, N. K., Kim, S. S., & Agarwal, J. (2004). Internet users' information privacy concerns (IUIPC): The construct, the scale, and a causal model. *Information Systems Research*, 15(4), 336–355.
34. Manser Payne, E. H., Peltier, J., & Barger, V. A. (2021). Enhancing the value co-creation process: Artificial intelligence and mobile banking service platforms. *Journal of Research in Interactive Marketing*, 15(1), 68–85.
35. Mariani, M. M., Perez-Vega, R., & Wirtz, J. (2022). AI in marketing, consumer research and psychology: A systematic literature review and research agenda. *Psychology & Marketing*, 39(4), 755–776.
36. Martin, K. D., Borah, A., & Palmatier, R. W. (2017). Data privacy: Effects on customer and firm performance. *Journal of Marketing*, 81(1), 36–58.
37. Martin, K. D., & Murphy, P. E. (2017). The role of data privacy in marketing. *Journal of the Academy of Marketing Science*, 45(2), 135–155.
38. Mayer, R. C., Davis, J. H., & Schoorman, F. D. (1995). An integrative model of organizational trust. *Academy of Management Review*, 20(3), 709–734.
39. McKnight, D. H., Choudhury, V., & Kacmar, C. (2002). Developing and validating trust measures for e-commerce: An integrative typology. *Information Systems Research*, 13(3), 334–359.
40. Nunnally, J. C. (1978). *Psychometric theory* (2nd ed.). McGraw-Hill.
41. Pavlou, P. A. (2003). Consumer acceptance of electronic commerce: Integrating trust and risk with the technology acceptance model. *International Journal of Electronic Commerce*, 7(3), 101–134.
42. Podsakoff, P. M., MacKenzie, S. B., Lee, J.-Y., & Podsakoff, N. P. (2003). Common method biases in behavioral research: A critical review of the literature and recommended remedies. *Journal of Applied Psychology*, 88(5), 879–903.
43. Puntoni, S., Reczek, R. W., Giesler, M., & Botti, S. (2021). Consumers and artificial intelligence: An experiential perspective. *Journal of Marketing*, 85(1), 131–151.
44. Rai, A. (2020). Explainable AI: From black box to glass box. *Journal of the Academy of Marketing Science*, 48(1), 137–141.
45. Shin, D. (2021). The effects of explainability and causability on perception, trust, and acceptance: Implications for explainable AI. *International Journal of Human-Computer Studies*, 146, 102551.

46. Shin, D. (2020). User perceptions of algorithmic decisions in the personalized AI system: Perceptual evaluation of fairness, accountability, transparency, and explainability. *Journal of Broadcasting & Electronic Media*, 64(4), 541–565.
47. Shin, D., & Park, Y. J. (2019). Role of fairness, accountability, and transparency in algorithmic affordance. *Computers in Human Behavior*, 98, 277–284.
48. Smith, H. J., Dinev, T., & Xu, H. (2011). Information privacy research: An interdisciplinary review. *MIS Quarterly*, 35(4), 989–1015.
49. Tucker, C. E. (2014). Social networks, personalized advertising, and privacy controls. *Journal of Marketing Research*, 51(5), 546–562.
50. Venkatesh, V., & Davis, F. D. (2000). A theoretical extension of the technology acceptance model: Four longitudinal field studies. *Management Science*, 46(2), 186–204.
51. 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.