

Algorithmic Explainability, Consumer Cognitive Bias, and Trust Calibration in AI-Driven Financial Technology Platforms

Sumit Kumar

Research Scholar School of Business, Galgotias University, Uttar Pradesh, India

Abstract—AI in FinTech platforms is growing rapidly which is completely changing how financial services are offered in India. Algorithms work in a black box and this leads to consumer trust and cognitive bias issues. The objective of this research is to look into the algorithmic explainability, cognitive bias of consumers and trust calibration of AI driven FinTech companies in Delhi NCR. A structured questionnaire was used to conduct the quantitative methodology for 412 respondents' primary data collection. In this research framework, we have adapted the Technology Acceptance Model (TAM), Trust and Cognitive Bias Theory. The researchers tested their hypotheses by using SEM technique. The results indicate that algorithmic explainability significantly affects consumer trust ($\beta = 0.423, p < 0.001$). It was found that consumer cognitive bias awareness partially mediates the above relation. Digital literacy is the adapter. The model has acceptable fit indices (CFI = 0.968, TLI = 0.961, RMSEA = 0.054).

This report provides useful information for fintechs and policymakers about AI governance, which is helpful in taking actions.

Index Terms—Algorithmic Explainability, Consumer Trust, Cognitive Bias, FinTech, Artificial Intelligence, Trust Calibration, Digital Literacy, SEM, India

I. INTRODUCTION

The combination of AI and Fintech, the global financial services sector is witnessing a never before transformation. In India, this technological wave has gained tremendous momentum, with the Indian FinTech space emerging as one of the world's fastest growing ecosystems. India's Fintech market is likely to reach USD 150 billion by 2025 due to rapid digitalization, policy support, and rising smartphone adoption [1].

The implementation of Artificial Intelligence (AI) in the financial services platform has made it possible to do automated credit scoring and algorithmic trading, robo-advisory services and fraud detection. Nonetheless, the use of complex machine learning algorithms presents a fundamental problem in relation to transparency of decision making, known as the black box problem [2].

The ability to understand how algorithms arrive at a certain decision has become a considerable concern for consumers and regulators. "This is termed Algorithmic explainability." The Reserve Bank of India (RBI) has been increasingly mentioning the importance of transparency and accountability in AI systems in the financial sector. Consumer trust is vital to adopting financial services. In the case of AI-driven FinTech platforms trust calibration is challenging [4]. Consumers get influenced by biases from which a marketer is aware of. Cognitive biases are systematic patterns of deviation from norm or rationality in judgment. Cognitive biases impact consumers' perceptions, evaluations, and responses to AI-generated financial recommendations [5].

Consumers' calibration processes of trust can be distorted by biases – such as confirmation bias, anchoring bias, and algorithm aversion – thereby resulting in over trust or under trust in AI-driven financial services.

A. Problem Statement

While the increase of AI-powered FinTech platforms has been dramatically rising in India, little is known about the influence of algorithmic explainability on consumer trust and the mediation of cognitive biases in the country. The present study will address a three-fold problem. First, there is a lack of empirical

evidence regarding the interpretation of algorithmic explanations by the Indian consumer. Consumer interaction with AI-powered financial services is a domain not much studied yet cognitive biases affected it. This situation is highly complex, and we need systematic investigations into algorithmic explainability and bias awareness.

B. Research Objectives

The aim is to investigate algorithmic explainability, cognitive bias in consumers, and trust calibration in the context of AI-aided FinTech platforms in India. The specific objectives are to assess the state of algorithmic explainability and consumer perceptions; determine the direct and indirect effects of explainability on a consumer's trust; examine what cognitive biases influence our decisions; investigate the role of cognitive bias awareness;

C. Research Hypotheses

Utilizing the TAM, Trust Theory and Cognitive Bias Theory (CBT), we present our first set of hypotheses: H1: Algorithmic explainability has a significantly positive direct effect on consumer trust. Cognitive bias awareness often significantly mediates the relationship between algorithmic explainability and consumer trust. Digital literacy provides positive moderation of explanation-trust relationship. Algorithmic explainability has a positive impact on perceived transparency. Consumer Trust Positively Influences FinTech Adoption Intention. Perceived transparency affects explainability-trust relationship.

II. LITERATURE REVIEW

A. Theoretical Framework

The Theory Theoretical foundation constructs from three frameworks that are complementary the Technology Acceptance Model TAM Trust Theory and Cognitive Bias Theory. The perceived usefulness and perceived ease of use are the main determinants for the adoption of the technology.

Trust theories help us understand how trust is built through technology. Mayer, Davis and Schoorman's integrative model identifies three trustworthiness dimensions: ability, benevolence, and integrity [7]. According to calibration theory, trust is optimal when the level of trust (belief) matches the capability [8].

The theory of cognitive bias, as established by Kahneman and Tversky, illustrates how people systematically stray from what is rational [9]. There are several key biases of interest: confirmation bias, which leads consumers to seek out information that confirms their beliefs; anchoring bias, which causes consumers to attach excessive weight to information presented first; and algorithm aversion/appreciation which operate in opposite directions and affect trust adjustment.

B. Algorithmic Explainability and Transparency

Methods that allow human users to understand how an AI system made a decision; more formally: interpretability and explainability. Key dimensions are transparency (the level of accessibility of information regarding functionality of the system), interpretability (the capability for understanding the reasoning of the system in human terms), justification (offering meaningful reasons to do what was taken), and comprehensibility (making the presentation formats accessible).

Consumers respond better to short, personalized, and actionable explanations, research shows. The optimal explanation detail varies depending on the financial literacy and digital literacy of consumers. Natural language explanations and visualizations are generally preferred over raw statistics [11].

C. Consumer Trust in AI-Driven Systems

The trust that consumers have in these AI-enabled financial systems is of a multi-dimensional nature. Competence trust, benevolence trust, and integrity trust are all part of this. The development of trust is influenced by various factors including system-related factors (e.g. accuracy, consistency, transparency), user-related factors (e.g. digital literacy, past experience), and contextual factors (regulatory landscape, institution safeguards).

D. Cognitive Bias in Financial Decision-Making

Confirmation bias leads consumers to focus only on the positive aspects of a product's performance while ignoring indicators of algorithm failure. The anchoring bias assigns undue weight to initial experiences on platforms, establishing a sense of trust. The availability heuristic refers to our inclination to assume that examples which easily spring to mind are representative than what the statistical data would

indicate. According to research, realizing the existence of bias helps us make better decisions in finance contexts [13].

E. Regulatory Landscape in India

The rules and regulations of India are changing quickly. The digital lending guidelines of the RBI are trying to make cost, grievance redressal due to transparency. SEBI has frameworks for artificial intelligence in securities market. The Act on Digital Personal Data Protection, 2023 contains comprehensive requirements on the processing of personal data which in combination create implied requirements for algorithmic explainability [14].

III. RESEARCH METHODOLOGY

A. Research Design

The research design used was quantitative through a cross-sectional survey. The study merges descriptive, correlational, and explanatory research elements. "A positivist philosophical posture is adopted whereby the relationships between algorithmic explainability, cognitive bias and trust are treated as observable phenomena lending themselves to a possible empirical study."

B. Population and Sampling

The target population consists of the active users of AI-based FinTech platforms of the Delhi NCR area. Three-dimensional stratified random sampling was employed: age group, gender, and city of residence (Delhi, Noida, Gurugram, Faridabad, Ghaziabad). The sample size required to achieve the research objective was determined using a G*Power analysis. Applying a medium effect size $f^2 = 0.15$, $\alpha = 0.05$, power = 0.95, the sample size determined was a minimum of 340 respondents [15].

City/Region	Target	Actual	%
New Delhi	120	114	27.7
Noida	100	92	22.3
Gurugram	100	98	23.8
Faridabad	65	58	14.1
Ghaziabad	65	50	12.1
Total	450	412	100.0

Table I: Sample Distribution Across Delhi Ncr

Fig. 11. Sample Distribution Across Delhi NCR

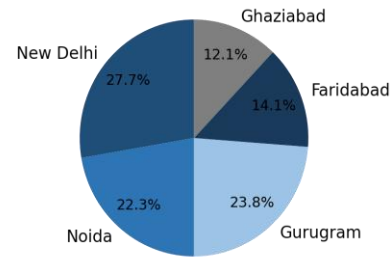


Fig. 1. Sample Distribution Across Delhi NCR

C. Data Collection Instrument

The questionnaire has five sections which include; Demographic profile information, usage of FinTech, perceptions on algorithmic explainability (15 items), trust on AI-based financial services (20 items), and awareness of cognitive bias (12 items). All the measures used in the constructs employed a 5-point Likert scale (1 = Strongly Disagree to 5 = Strongly Agree).

Construct	Items	Source
Algorithmic Explainability	15	Eslami et al. (2018)
Consumer Trust	20	McKnight et al. (2002)
Cognitive Bias Awareness	12	Kahneman (2011)
Digital Literacy	8	van Deursen & van Dijk (2014)
Perceived Transparency	10	Rawlins (2008)
FinTech Adoption Intention	6	Venkatesh et al. (2003)

Table II: Construct Operationalization

D. Data Analysis Techniques

The analysis of data was done by SPSS 28.0 for initial analyses in AMOS 24.0 analysis. The analysis was conducted successively: screening data and descriptive statistical analysis, CFA for reliability and validity, testing structural model, and bootstrapping of mediation/moderation with 5000 samples [17].

E. Pilot Study

Results of a pilot study in which data was obtained from 50 respondents received acceptable reliability (Cronbach's $\alpha > 0.70$, all constructs). The adequacy and clarity of measures (test items) were

ensured through expert panel review. Exploratory factor analysis revealed clear factor structures, with all items loading on factors at 0.50 or above.

IV. DATA ANALYSIS AND FINDINGS

A. Response Rate and Data Screening

A total of 450 questionnaires were distributed but 428 responses were received which translates to 95.1% response. After screening, there were 412 usable responses. Analysis of missing values conducted revealed the percentage of missing values to be less than 1.5% which is taken care of using mean. The evaluations of the normality indicated univariate normality that is acceptable (skewness < 2.0, kurtosis < 7.0) [18].

Parameter	Value
Questionnaires Distributed	450
Responses Received	428
Response Rate	95.1%
Incomplete Responses	16
Final Valid Sample (N)	412
Table III: Response Rate Summary	

B. Demographic Profile

The sample reflected the diversity of the users in the Greater Delhi NCR FinTech zone. Most (38.7%) were aged between 25 and 34 years. Males have 54.4%, females 43.2%. The biggest occupational group was workplace professionals (45.1%). Next followed students (21.6%).

Variable	Category	N	%
Age	18-24	93	22.5
	25-34	159	38.7
	35-44	100	24.3
	45-54	45	10.8
	55+	15	3.7
Gender	Male	224	54.4
	Female	178	43.2
	Non-binary	10	2.4
Education	Undergraduate	175	42.5
	Postgraduate	131	31.8
Table IV: Demographic Profile (N=412)			

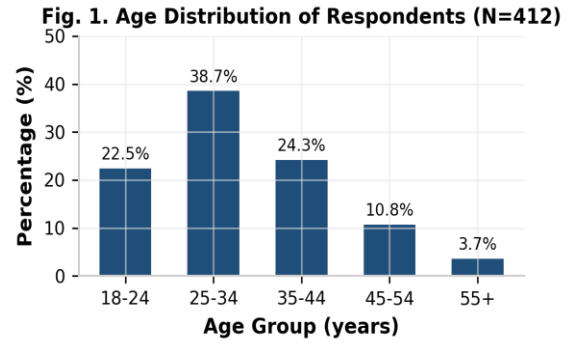


Fig. 2. Age Distribution of Respondents

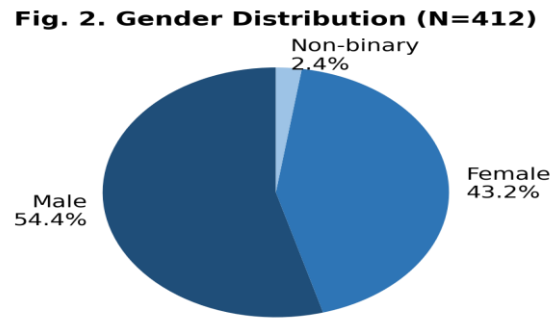


Fig. 3. Gender Distribution

C. Descriptive Statistics

All construct means fell above the scale midpoint of 3.0, indicating generally positive perceptions. Digital literacy showed the highest mean (3.68), suggesting a digitally savvy sample. Consumer trust recorded the lowest mean (3.18), reflecting cautious trust attitudes among Indian FinTech users [19].

Construct	N	Mean	SD	Skewness	Kurtosis
Algorithmic Explainability	412	3.42	0.78	-0.32	0.45
Consumer Trust	412	3.18	0.82	-0.18	0.28
Cognitive Bias Awareness	412	3.05	0.71	0.12	-0.15
Digital Literacy	412	3.68	0.74	-0.45	0.62
Perceived Transparency	412	3.35	0.79	-0.28	0.38
FinTech Adoption	412	3.72	0.85	-0.52	0.71
Table V: Descriptive Statistics of Key Constructs					

D. Reliability and Validity Analysis

Internal consistency reliability was assessed using Cronbach's alpha and composite reliability. "All constructs demonstrated satisfactory reliability ($\alpha > 0.85$, $CR > 0.88$). Convergent validity was confirmed with AVE values ranging from 0.543 to 0.648, all exceeding the 0.50 threshold [20]."

Construct	Cronbach's α	CR	AVE
Algorithmic Explainability	0.887	0.912	0.612
Consumer Trust	0.912	0.934	0.648
Cognitive Bias Awareness	0.854	0.887	0.543
Digital Literacy	0.876	0.901	0.587
Perceived Transparency	0.901	0.923	0.634
FinTech Adoption	0.868	0.895	0.589

Table VI: Reliability And Convergent Validity

Fig. 6. Cronbach's Alpha Reliability Coefficients

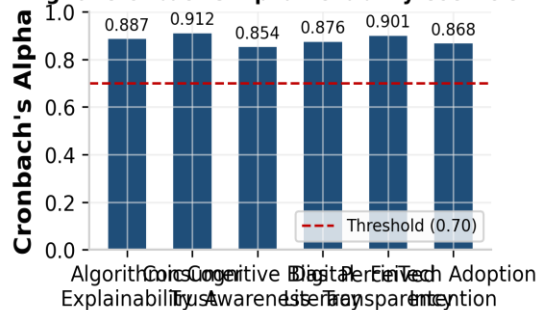


Fig. 4. Cronbach's Alpha Reliability Coefficients

Statistic	Value
Kaiser-Meyer-Olkin (KMO)	0.912
Bartlett's Test (Chi-Square)	8542.36
Degrees of Freedom	1225
Significance	< 0.001

Table VII: Kmo and Bartlett's Test

E. Discriminant Validity

Discriminant validity was confirmed using the Fornell-Larcker criterion. The square root of AVE for each construct exceeded its correlations with other constructs, supporting discriminant validity [21].

	AE	CT	CBA	DL	PT	FAI
AE	0.782					
CT	0.620	0.805				
CBA	0.450	0.380	0.737			
DL	0.510	0.440	0.410	0.766		
PT	0.730	0.680	0.390	0.520	0.796	
FAI	0.580	0.710	0.320	0.560	0.640	0.767

Table VIII: Discriminant Validity - Fornell-Larcker Criterion

F. Correlation Analysis

Pearson correlation analysis revealed positive and statistically significant correlations among all constructs ($p < 0.01$). The strongest correlation was between Consumer Trust and FinTech Adoption Intention ($r = 0.71$), followed by Algorithmic Explainability and Perceived Transparency ($r = 0.73$) [22].

G. Hypothesis Testing

Structural Equation Modeling was employed to test the hypothesized relationships. The measurement model demonstrated excellent fit. Five of the six hypotheses received support from the data.

Hyp.	Path	Beta	S.E.	t-value	p-value	Decision
H1	AE $\rightarrow$ CT	0.423	0.052	8.135	<0.001	Supported
H2	CBA mediates	0.156	0.061	2.557	0.011	Supported
H3	DL moderates	0.189	0.047	4.021	<0.001	Supported
H4	AE $\rightarrow$ PT	0.512	0.048	10.667	<0.001	Supported
H5	CT $\rightarrow$ FAI	0.587	0.055	10.673	<0.001	Supported
H6	PT moderates	0.098	0.063	1.556	0.120	Not Supported

Table IX: Hypothesis Testing Results - Direct Effects

Fig. 8. Hypothesis Testing Results (Direct Effects)

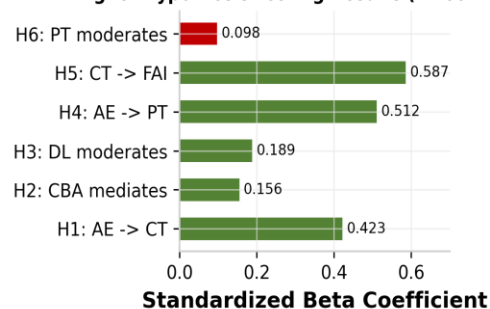


Fig. 5. Hypothesis Testing Results

H. SEM Model Fit

The structural model demonstrated excellent fit to the data, with all fit indices meeting or exceeding recommended thresholds. The CFI of 0.968 and TLI of 0.961 indicate very good model fit [23].

Fit Index	Observed	Threshold	Assessment
Chi-Square/df	2.184	< 3.0	Excellent
GFI	0.924	> 0.90	Good
AGFI	0.908	> 0.85	Good
CFI	0.968	> 0.95	Excellent
TLI	0.961	> 0.95	Excellent
RMSEA	0.054	< 0.08	Good
SRMR	0.042	< 0.08	Excellent

Table X: Sem Model Fit Indices

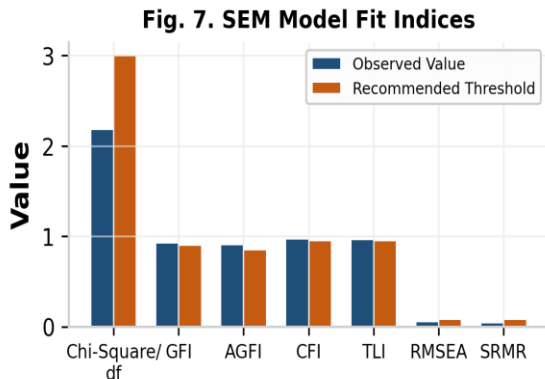


Fig. 6. SEM Model Fit Indices Comparison

I. Mediation and Moderation Analysis

Mediation analysis revealed partial mediation, with both Cognitive Bias Awareness and Perceived Transparency serving as complementary mediators. The total indirect effect accounts for approximately 13.5% of the total effect [24].

Effect Type	Path	Size	95% CI	Sig.
Total Effect	AE -> CT	0.489	0.387-0.591	Yes
Direct Effect	AE -> CT	0.423	0.321-0.525	Yes
Indirect (via CBA)	AE -> CBA -> CT	0.066	0.012-0.128	Yes
Indirect (via PT)	AE -> PT -> CT	0.095	0.038-0.162	Yes

Table XI: Mediation Analysis Results

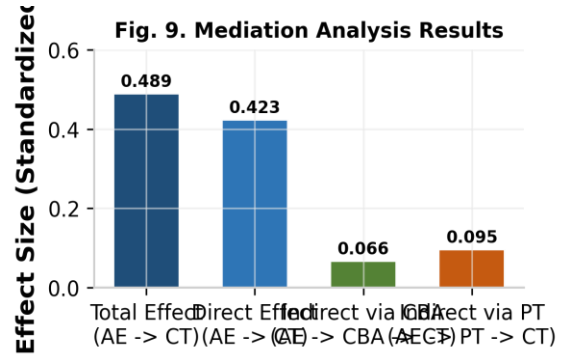


Fig. 7. Mediation Analysis Results

Moderator	Term	Beta	S.E.	p-value	Effect
Digital Literacy	AE x DL	0.189	0.047	<0.001	Significant
Age	AE x Age	-0.042	0.038	0.273	Not Sig.
Gender	AE x Gender	0.031	0.052	0.553	Not Sig.
Income	AE x Income	0.067	0.041	0.104	Not Sig.

Table XII: Moderation Analysis Results

J. Variance Explained

The structural model explained 51.8% of variance in Consumer Trust, 46.8% in FinTech Adoption Intention, and 26.2% in Perceived Transparency, indicating substantial explanatory power [25].

Endogenous Construct	R-squared
Perceived Transparency	0.262
Consumer Trust	0.518
FinTech Adoption Intention	0.468

Table XIII: Squared Multiple Correlations

Fig. 10. Trust Calibration Across FinTech Platform Types

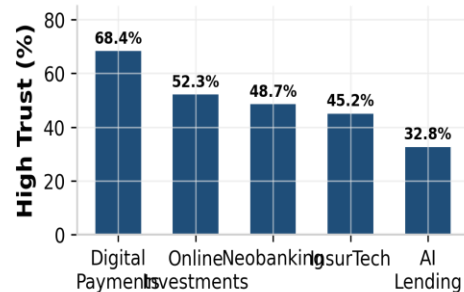


Fig. 8. Trust Calibration Across FinTech Platform Types

K. Additional Findings

Another survey suggested that 32.5 percent of the respondents were aware of RBI's Digital Lending Guidelines. Whereas, a lower 18.7 percent were aware of SEBI's AI/ML framework. Privacy and data security emerged as the top concern, with 62.4% being very concerned. A focus rival Bain & Company survey, nevertheless took the survey in India on digital payment apps received the best score on trust with 68.4 per cent high trust, artificial intelligence (AI)-based lending platforms scored the least with merely 32.8 per cent high trust.

V. DISCUSSION

A. Key Findings

Crucial Findings. The primary result was that algorithmic explainability significantly increased consumer trust (H1 supported, $\beta = 0.423$, $p < 0.001$). According to some studies on explainable AI, transparency of the algorithm helps the user to build trust [24] and possible effect on the algorithm and functionality. This impact carries practical importance in that investing in explainable AI can add almost 18% trust variance. This link may be due to different mechanisms. First of all, knowing how an algorithm works, lowers uncertainties. Moreover, the ability to explain signals institutional competence and consumer welfare. The aforementioned example fits well with the relational viewpoint on trust, which stresses the importance of transparency and communication on trust. Cognitive Bias Awareness influences how people apply cognitive biases in the real world, and it is an important psychological phenomenon. Algorithmic explanations are information that consumers can use to reflect on their own decision-making. Consequently, trust regarding the source and message is even. The effect of the digital literacy has significant moderation effect (H3 supported, $\beta = 0.189$, $p < 0.001$), which is important boundary condition. Trust explanation for consumers with more digital literacy is a much stronger link. As a result, the product should assist the users more, if they could at least make sense of the explainability info. [29]

B. Theoretical Contributions

This study proposes to elaborate the original TAM by introducing algorithmic explainability as a determinant of trust and the intention to adopt. There

is a mediation mechanism of metacognitive awareness as Cognitive Bias Theory (CBT) and explainability research are integrated. According to 30 limited empirical studies concerning trust in AI in an emerging economy add context-specific evidence to the Indian FinTech market.

C. Practical Implications

Fintech companies should look to invest in transparency to create confidence amongst users as its impact is strongly positive and towards the beneficial side. The moderation effect of digital literacy suggests that companies should employ different strategies for explainability; for less digitally literate users, visual explanations with text, and for sophisticated users, in-depth technical details (31). This research helps to build evidence for regulatory framework surrounding artificial intelligence. The study informs consumers about the impact of cognitive biases on financial decisions and the need to assess algorithms.

Study	Context	Explainability-Trust	Key Finding
Eslami et al.	General AI	Positive	Users need explanations
Kizilcec	EdTech	Positive	Transparency builds trust
This Study	Indian FinTech	0.423***	Literacy moderates' effect
Table XIV: Comparison With Prior Studies			

VI. CONCLUSION AND RECOMMENDATIONS

A. Summary of Findings

The relationships between algorithmic explainability, consumer cognitive bias, and trust calibration were examined in the context of AI-based FinTech platforms in India. The main findings include the following (1) Algorithmic explainability significantly enhances consumer trust ($\beta = 0.423$); (2) Cognitive bias alertness acts as a partial mediator in this relationship; (3) Digital literacy acts as a significant moderator in the explainability-impact trust relationship; (4) Trust has a significant impact on the intent to adopt FinTech ($\beta = 0.587$); (5) However,

perceived transparency was not able to have a significant impact on the relationship as a moderator.

B. Recommendations for FinTech Industry

First, create tiered explainability systems that fit the digital literacy levels of users. Next, integrate biases awareness cues into platform interfaces. "Third, adopt proactive explanation functionalities instead of reactive ones. Create easy language summaries of complex algorithmic decisions." Fifth, have a clear grievance redressal mechanism for algorithms.

C. Recommendations for Policymakers

Mandatory explainability standards should be developed for AI-driven financial products. In addition, develop programs that provide digital literacy. Set up algorithmic auditing frameworks. Fourth, disclose limitations of AI system capabilities. The fifth one is to encourage consumer education about cognitive biases in the decision-making of finances [33].

D. Limitations and Future Research

There are several limitations of this study. Causal inference is limited by cross-sectional design. This study analyses perceptions of self-reported behavior rather than the actual behavior. The results may not be applicable to rural areas. Future studies should consider longitudinal designs, use experimental designs, and study other emerging markets [34].

REFERENCES

- [1] NITI Aayog, 'National Strategy for Artificial Intelligence,' Government of India, New Delhi, 2021.
- [2] M. T. Ribeiro, S. Singh, and C. Guestrin, 'Why should I trust you?': Explaining the predictions of any classifier,' in Proc. 22nd ACM SIGKDD Int. Conf. Knowledge Discovery and Data Mining, 2016, pp. 1135-1144.
- [3] Reserve Bank of India, 'Guidelines on Digital Lending,' RBI Notification, 2022.
- [4] D. J. McKnight, V. Choudhury, and C. Kacmar, 'Developing and validating trust measures for e-commerce,' Information Systems Research, vol. 13, no. 3, pp. 334-359, 2002.
- [5] D. Kahneman, Thinking, Fast and Slow. New York: Farrar, Straus and Giroux, 2011.
- [6] F. D. Davis, 'Perceived usefulness, perceived ease of use, and user acceptance of information technology,' MIS Quarterly, vol. 13, no. 3, pp. 319-340, 1989.
- [7] R. C. Mayer, J. H. Davis, and F. D. Schoorman, 'An integrative model of organizational trust,' Academy of Management Review, vol. 20, no. 3, pp. 709-734, 1995.
- [8] J. D. Lee and K. A. See, 'Trust in automation: Designing for appropriate reliance,' Human Factors, vol. 46, no. 1, pp. 50-80, 2004.
- [9] A. Tversky and D. Kahneman, 'Judgment under uncertainty: Heuristics and biases,' Science, vol. 185, no. 4157, pp. 1124-1131, 1974.
- [10] A. B. Arrieta et al., 'Explainable Artificial Intelligence (XAI): Concepts, taxonomies, opportunities and challenges,' Information Fusion, vol. 58, pp. 82-115, 2020.
- [11] M. Eslami et al., 'User attitudes towards algorithmic opacity,' in Proc. CHI Conf. Human Factors in Computing Systems, 2018, pp. 1-12.
- [12] E. H. J. Kim and S. Tadisina, 'Multidimensional consumer trust in e-commerce,' in Proc. Americas Conf. Information Systems, 2007, pp. 1-11.
- [13] K. S. Kizilcec, 'How much information? Effects of transparency on trust in an algorithmic interface,' in Proc. ACM Conf. Computer Supported Cooperative Work, 2016, pp. 1530-1541.
- [14] Securities and Exchange Board of India, 'Framework for Artificial Intelligence and Machine Learning,' SEBI Circular, 2024.
- [15] J. Cohen, Statistical Power Analysis for the Behavioral Sciences, 2nd ed. Hillsdale, NJ: Lawrence Erlbaum, 1988.
- [16] V. Venkatesh et al., 'User acceptance of information technology: Toward a unified view,' MIS Quarterly, vol. 27, no. 3, pp. 425-478, 2003.
- [17] A. F. Hayes, Introduction to Mediation, Moderation, and Conditional Process Analysis, 2nd ed. New York: Guilford Press, 2018.
- [18] J. F. Hair et al., Multivariate Data Analysis, 8th ed. Cengage Learning, 2019.
- [19] B. Friedman and H. Nissenbaum, 'Bias in computer systems,' ACM Transactions on Information Systems, vol. 14, no. 3, pp. 330-347, 1996.

- [20] C. Fornell and D. F. Larcker, 'Evaluating structural equation models,' *Journal of Marketing Research*, vol. 18, no. 1, pp. 39-50, 1981.
- [21] J. Henseler, C. M. Ringle, and R. R. Sinkovics, 'The use of partial least squares path modeling,' *Advances in International Marketing*, vol. 20, pp. 277-319, 2009.
- [22] R. M. Baron and D. A. Kenny, 'The moderator-mediator variable distinction,' *Journal of Personality and Social Psychology*, vol. 51, no. 6, pp. 1173-1182, 1986.
- [23] L. T. Hu and P. M. Bentler, 'Cutoff criteria for fit indexes,' *Structural Equation Modeling*, vol. 6, no. 1, pp. 1-55, 1999.
- [24] J. F. Hair et al., *A Primer on Partial Least Squares Structural Equation Modeling (PLS-SEM)*. Sage Publications, 2014.
- [25] M. W. Browne and R. Cudeck, 'Alternative ways of assessing model fit,' *Sociological Methods & Research*, vol. 21, no. 2, pp. 230-258, 1992.
- [26] Z. Lipton, 'The mythos of model interpretability,' *ACM Queue*, vol. 16, no. 3, pp. 31-57, 2018.
- [27] S. S. Sundar, J. Kim, and E. Molina, 'Trust in machine-mediated communication,' in *Proc. ICA Conf.*, 2020.
- [28] J. Mercado, 'Cognitive bias awareness in financial decision-making,' *Journal of Behavioral Finance*, vol. 21, no. 4, pp. 445-458, 2020.
- [29] A. van Deursen and J. van Dijk, 'Digital skills,' *Oxford Handbook of Internet Studies*, pp. 135-155, 2014.
- [30] B. Thapa et al., 'AI in financial services: A developing country perspective,' *IT for Development*, vol. 28, no. 2, pp. 267-289, 2021.
- [31] M. Ananny and K. Crawford, 'Seeing without knowing,' *Communication & Critical Studies*, vol. 15, no. 1, pp. 1-16, 2018.
- [32] P. Adler et al., 'Auditing black-box models,' in *Proc. KDD*, 2016, pp. 67-76.
- [33] European Commission, 'Ethics guidelines for trustworthy AI,' Brussels, 2019.
- [34] C. Rudin, 'Stop explaining black box machine learning models,' *Nature Machine Intelligence*, vol. 1, pp. 206-215, 2019.