

Original Article

AI-Enabled Financial Advice and Investment Decision Quality: A Dual-Moderated Mediation Framework for Indian Retail Investors

Dr. Shubham¹, Dr. Suchitra Ranglani²

¹ Assistant Professor, Government Degree College, Faridpur, Bareilly, Uttar Pradesh, India

² Former Assistant Professor, School of Economics and Commerce, KIIT Deemed to be University, Bhubaneswar, India

Manuscript ID:
BN-2025-0204016

ISSN: 3065-7865

Volume 2

Issue 4

April 2025

Pp 89-99

Submitted: 31 Jan 2025

Revised: 06 Feb 2025

Accepted: 11 Mar 2025

Published: 30 Apr 2025

DOI:
10.5281/zenodo.22122690

DOI link:
<https://doi.org/10.5281/zenodo.22122690>



Quick Response Code:



Website: <https://bnir.us>



Abstract

AI-enabled financial advice can lower the cost of analysis and personalization, yet adoption does not establish that retail investors understand a recommendation or use it appropriately. This conceptual paper develops a dual-moderated mediation framework for explaining when AI-enabled advice improves investment decision quality among Indian retail investors. A problem-driven synthesis integrates financial-capability research, information-adoption theory, human-automation trust, and recent robo-advisory evidence with official Indian market and regulatory sources. The model treats AI-advice design quality suitability, explanation transparency, uncertainty communication, and data provenance as the antecedent. Advice comprehension and verification are specified as the task-specific mediator, while baseline digital financial literacy moderates the first-stage relationship. Trust is separated from observed reliance, and trust-reliance calibration is proposed as the second-stage boundary condition. The outcome is criterion-scored decision quality rather than adoption intention or short-term return. Six propositions predict that design quality improves comprehension and verification, which in turn improves decision quality; the indirect effect should be strongest when baseline digital financial literacy is high and reliance is calibrated to demonstrated system capability. India is not treated merely as a location. SEBI's national survey reveals a large awareness-participation gap, marked urban-rural heterogeneity, and strong dependence on social and influencer information, making verification capability and calibrated reliance theoretically consequential. The paper contributes construct clarity, a testable model, a multi-stage empirical design, and platform and policy implications for responsible AI-enabled investment advice.

Keywords: artificial intelligence; robo-advice; digital financial literacy; advice comprehension; verification behavior; calibrated trust; investment decision quality; India

Introduction

Artificial intelligence (AI) now enters retail investing through automated portfolio allocation, risk profiling, product screening, brokerage nudges, and conversational interfaces. These systems can process large information sets, tailor outputs to stated preferences, and expand access to standardized guidance. Recent reviews nevertheless show that the design, business model, and empirical evaluation of robo-advisory remain fragmented, while generative systems add new risks concerning provenance, hallucination, and persuasive fluency (Cardillo & Chiappini, 2024; OECD, 2025; Zhu et al., 2024). The relevant scholarly question is therefore no longer whether investors will adopt an AI tool, but when an AI-mediated recommendation improves the quality of a consequential financial choice. Adoption and decision quality are analytically different. Technology-acceptance research explains why a system may appear useful or easy to use (Davis, 1989), and financial robo-advice studies have appropriately examined adoption and acceptance (Belanche et al., 2019; Lourenço et al., 2020). A favorable adoption intention, however, does not reveal whether the user identified an unsuitable horizon, checked a fee, noticed stale data, or resisted an implausibly precise forecast. A platform can increase engagement while worsening welfare if it makes action easier than evaluation. Conversely, a cautious investor may reject useful automation after observing a single error.

Creative Commons (CC BY-NC-SA 4.0)

This is an open access journal, and articles are distributed under the terms of the [Creative Commons Attribution-NonCommercial-ShareAlike 4.0 International](https://creativecommons.org/licenses/by-nc-sa/4.0/) Public License, which allows others to remix, tweak, and build upon the work noncommercially, as long as appropriate credit is given and the new creations are licensed under the identical terms.

Address for correspondence:

Dr. Shubham Assistant Professor, Government Degree College, Faridpur, Bareilly, Uttar Pradesh, India
Email: shubhisinha23@gmail.com

How to cite this article:

Shubham, & Ranglani, S. (2025). AI-Enabled Financial Advice and Investment Decision Quality: A Dual-Moderated Mediation Framework for Indian Retail Investors. *Bulletin of Nexus Journal*, 2(4), 89–99.
<https://doi.org/10.5281/zenodo.22122690>

Evidence on algorithm aversion and algorithm appreciation shows that machine advice can be either discounted or overweighted, depending on the task, comparison standard, and user experience (Dietvorst et al., 2015, 2018; Logg et al., 2019). Digital financial literacy (DFL) and trust are central variables, but their causal ordering is often underspecified. Baseline DFL is a relatively durable capability that exists before a focal advice episode; it cannot plausibly be caused by a single exposure to AI advice and then serve as a mediator in the same cross-sectional model. Trust, reliance, and calibration are also not interchangeable. Trust is an attitude under vulnerability, reliance is observable behavior, and calibration is the degree to which confidence and reliance correspond to actual system capability (Hoff & Bashir, 2015; Lee & See, 2004). Treating all three as one self-report scale obscures both underuse and overuse. This paper resolves those problems by proposing a dual-moderated mediation model. AI-advice design quality is the antecedent; advice comprehension and verification are the task-specific mediator; baseline DFL moderates the first stage; and trust-reliance calibration conditions the conversion of competent evaluation into decision quality. Investment decision quality is assessed against scenario-specific criteria such as suitability, diversification, cost sensitivity, and justified override behavior, not by short-term return. The model is developed for Indian retail investors because the Indian context combines rapid digital access with sharp differences in awareness, participation, education, language, and information sources. The paper makes four contributions. First, it replaces an adoption-centered account with a decision-quality account. Second, it separates stable capability from episode-specific cognitive and verification processes. Third, it distinguishes attitudinal trust, behavioral reliance, and calibration, thereby making the risk of overreliance

empirically visible. Fourth, it turns India into a theoretical boundary context by linking national evidence and the regulatory perimeter to predicted heterogeneity in the model. The remainder of the paper defines the scope and conceptual approach, develops the constructs and propositions, explains the Indian boundary conditions, and specifies an empirical design and implications.

Scope and Conceptual Development Approach

1. Scope of AI-enabled financial advice

AI-enabled financial advice is defined here as retail-investor-facing guidance in which a computational system materially generates, personalizes, ranks, or explains a recommendation. The unit of analysis is an advice episode: a specific investor receives an output for a specific decision under identifiable system conditions. This definition excludes purely administrative automation and market information that is neither ranked nor connected to a decision. It includes both regulated advice and decision support, but it does not assume that every AI-generated statement is legally “investment advice.” That determination depends on personalization, product specificity, business conduct, and the applicable regulatory regime.

The category is heterogeneous. A rules-based robo-adviser may translate a structured risk profile into a model portfolio; a machine-learning system may screen products or forecast risk; and a general-purpose conversational model may explain, compare, or generate security-specific suggestions. These systems differ in error modes and accountability. The framework therefore applies to their common recommendation episode, while empirical studies should report the system class, provider, data sources, personalization level, and regulatory status rather than treating “AI” as a uniform treatment.

Table 1. Scope of retail-facing AI-enabled financial advice

System class	Decision role	Principal error or conduct risk	Boundary for this framework
Rules-based or hybrid robo-adviser	Risk profiling, model portfolio, rebalancing	Mis-specified inputs, unsuitable mapping, hidden fees or conflicts	Included when the output recommends or materially ranks an investor action
Machine-learning decision support	Screening, risk estimates, alerts, portfolio analytics	Distribution shift, opaque features, stale or biased data	Included when analytics are presented as actionable guidance
Conversational or generative AI	Explanation, comparison, question answering, generated suggestions	Hallucination, weak provenance, persuasive certainty, inconsistent outputs	Included only for a defined recommendation episode; not presumed to be regulated advice

2. Conceptual synthesis strategy

The paper is a conceptual theory-synthesis rather than a systematic review or an empirical study. Following problem-driven guidance for literature reviews (Snyder, 2019), it uses recent systematic and scoping reviews to map robo-advisory research, then integrates foundational work on information adoption, financial capability, and human-automation reliance. Peer-reviewed studies were prioritized for theoretical and empirical claims. Official SEBI, RBI, OECD, and IOSCO materials were used for measurement, market structure, and regulatory context. The literature and regulatory cutoff for the synthesis is 15 August 2024. The synthesis followed four analytic decisions. First, constructs were assigned a single causal role based on temporal ordering. Second, system exposure, user capability, episode-specific processing, and behavior were separated. Third, outcome validity was prioritized over convenient intention measures. Fourth, the Indian context was used to derive boundary conditions rather than appended as a descriptive paragraph. The review does not claim exhaustive database coverage, pooled effect sizes, or formal risk-of-bias assessment. Recent reviews themselves document that the empirical literature is dominated by adoption studies and that the welfare consequences of financial robo-advice remain underexamined (Cardillo & Chiappini, 2024; Nourallah et al., 2026; Wagner, 2024).

Construct Architecture

1. AI-advice design quality

AI-advice design quality is the extent to which an interface provides accurate, suitable, intelligible, uncertainty-aware, and verifiable guidance for the focal decision. It is not equivalent to a user's global perception that a system is "good." A rigorous study should manipulate or independently code design features and then measure perceptions as a proximal response. This distinction reduces common-method bias and allows researchers to identify which feature changed cognition. Four feature families are central. Suitability connects the recommendation to risk capacity, horizon, liquidity needs, existing holdings, and stated objectives. Explanation transparency makes the reasons and relevant inputs visible without pretending that every technical detail is meaningful to the user. Uncertainty communication reports confidence, ranges, limitations, and conditions under which the recommendation may fail. Provenance and recency identify the origin and date of data, distinguish generated content from verified records, and expose conflicts or commercial incentives. Contestability—

allowing users to correct inputs, compare alternatives, and seek human review—supports all four features. User research on AI-empowered financial advice indicates that nominal transparency can still feel incomprehensible, reinforcing the need to test actual understanding rather than disclosure presence alone (Zhu et al., 2023). Advice quality also requires a valid comparator. A fluent explanation cannot repair a substantively unsuitable recommendation, and "accuracy" cannot be inferred from past return alone. The experimental benchmark should specify what information was available, what recommendation met the task criteria, and how system reliability varied across trials. Systematic design work likewise emphasizes requirements extending beyond predictive performance (Namyslo et al., 2025). Studies comparing ChatGPT-type systems and robo-advisers show why outputs must be evaluated against an explicit decision standard rather than rhetorical quality (Oehler & Horn, 2024).

2. Baseline digital financial literacy as capability

DFL extends financial knowledge into digitally mediated environments. It combines understanding of risk, return, diversification, inflation, compounding, fees, and suitability with digital skills, knowledge of consumer rights, recognition of digital-finance risk, source verification, and protective behavior. Financial literacy is economically consequential for household choice (Lusardi & Mitchell, 2014). OECD/INFE measurement guidance treats financial knowledge, behavior, and attitudes as distinct, while Indian scale-development studies show that DFL is multidimensional rather than a synonym for app use (Chhillar et al., 2024; OECD, 2022, 2023; Ravikumar et al., 2022). In the proposed model, baseline DFL is measured before exposure to the focal system. This preserves causal ordering. An investor with stronger baseline DFL should be better able to interpret probabilities, detect an omitted fee, recognize a mismatch between risk tolerance and product volatility, and interrogate a source. Yet self-confidence must be measured separately: unjustified confidence can increase adoption without corresponding competence (Piehlmaier, 2022). Objective items, scenario tasks, and digital-risk questions should therefore be combined, while self-rated knowledge is retained as a distinct calibration variable. Baseline DFL is modeled as a first-stage moderator because capability changes what the investor can extract from a given design. A transparent explanation offers usable cues, but users need some knowledge to connect those cues to financial consequences. The

positive interaction proposed below is a parsimonious starting point. A compensatory alternative is also possible: structured explanations may benefit low-DFL users more. That rival should be preregistered and evaluated, particularly where interfaces use plain language, local examples, or multilingual support.

3. Advice comprehension and verification as the mechanism

Advice comprehension and verification are episode-specific processes. Comprehension means forming an accurate mental representation of what the system recommends, why it recommends it, what assumptions it uses, and what uncertainty remains. Verification means checking material claims, inputs, source status, fees, conflicts, and consistency with the investor's constraints before acting. The two processes are related but not identical: an investor may understand a recommendation and still fail to verify it, or may search for a second source without understanding why the sources differ. This mechanism is consistent with information-adoption theory, which links argument quality and source credibility to perceived usefulness and information use (Sussman & Siegal, 2003), but it is more demanding than perceived usefulness. In a financial task, the mediator should be demonstrated. Suitable measures include explanation-recall questions, numeracy-based interpretation, identification of a planted inconsistency, requests for source details, fee comparison, and the ability to state when the advice should not be followed. Cognitive forcing functions can improve deliberation in AI-supported decisions, although their benefits depend on effort and task design (Bućinca et al., 2021). The distinction also clarifies why an AI interface could influence the mediator without creating durable literacy. A well-designed episode can scaffold attention and verification today; repeated episodes may eventually produce learning, but that is a longitudinal claim. The immediate mediator is enacted comprehension and verification, while change in DFL is a separate longer-term outcome.

4. Trust, reliance, and calibration

Trust is an attitude that the system will help achieve the user's goals in a situation involving uncertainty and vulnerability. Reliance is a behavioral choice: accepting, modifying, rejecting, or seeking an alternative to the recommendation. Calibration is the correspondence between trust or reliance and the system's demonstrated capability for the relevant task. These distinctions follow established human-automation research (Hoff & Bashir, 2015; Lee & See, 2004) and are essential in finance, where both misplaced acceptance and

unjustified rejection can be costly. Trust has performance, process, and purpose components. Performance concerns competence and consistency; process concerns predictability, explanations, privacy, and data handling; purpose concerns incentives, benevolence, and alignment with the investor. Regulatory registration or a known brand may provide institutional assurance, but neither proves that a specific model is reliable for a specific task. Likewise, interface fluency can raise perceived competence without raising validity. Reviews of trust in AI show that human-like presentation and transparency can affect trust, but their effects depend on task and evidence (Glikson & Woolley, 2020). Calibration should not be measured by declaring high trust desirable. Across repeated trials with known system reliability, researchers can estimate appropriate acceptance, appropriate rejection, overreliance, and underreliance. A calibration score can compare subjective confidence with observed validity, while behavioral analyses can model whether participants discriminate between reliable and unreliable advice. This approach accommodates algorithm aversion after visible errors and algorithm appreciation when machine advice appears objective (Dietvorst et al., 2015; Logg et al., 2019). Limited user control may reduce aversion, but control should improve diagnosis rather than permit arbitrary tailoring of an answer (Dietvorst et al., 2018).

5. Investment decision quality

Investment decision quality is the degree to which a choice is suitable, informed, internally consistent, cost-aware, and supported by a defensible process. It is not synonymous with trading, acceptance, confidence, or realized return. Short-horizon returns are noisy and can reward a poor gamble; a diversified and suitable portfolio can experience a temporary loss. Behavioral intention is also insufficient because intention and action may diverge (Ajzen, 1991). The preferred outcome is a criterion score for standardized scenarios. Depending on the task, the scoring rubric can reward horizon and liquidity fit, diversification, avoidance of dominated or excessively costly products, correct interpretation of downside risk, attention to taxes and fees, and justified acceptance or override of advice. Outcomes should be blinded to condition and, where possible, scored by preregistered rules. Secondary outcomes may include information search, switching, concentration, turnover, and longitudinal adherence, but each must be interpreted against investor objectives rather than assumed to be good or bad.

Table 2. Constructs, causal roles, and defensible operationalization

Construct	Causal role	Definition	Preferred operationalization
AI-advice design quality	Antecedent (X)	Suitability, explanation transparency, uncertainty communication, and provenance/recency	Randomized interface features plus independent expert coding; perceptions measured separately
Baseline DFL	First-stage moderator (W_1)	Pre-existing financial and digital capability relevant to safe use	Pre-exposure objective knowledge, applied digital-risk tasks, and source-recognition items
Advice comprehension and verification	Mediator (M)	Accurate understanding and checking during the focal episode	Recall and interpretation items, planted-error detection, source and fee checks, information search
Attitudinal trust	Antecedent to reliance	Belief in competence, predictability, integrity, and goal alignment	Task-specific multi-item scale measured before the consequential choice
Trust–reliance calibration	Second-stage boundary condition (W_2)	Correspondence of confidence and reliance with demonstrated system capability	Trial-level appropriate acceptance/rejection and confidence–validity mismatch
Investment decision quality	Outcome (Y)	Suitable, informed, consistent, and cost-aware choice process	Blinded, preregistered scoring of portfolio scenarios; optional longitudinal process indicators

Proposed Dual-Moderated Mediation Framework

Figure 1 presents the proposed framework. AI-advice design quality should increase investment decision quality primarily by improving advice comprehension and verification. Baseline DFL conditions the first-stage effect because users differ in the capabilities needed to interpret design cues. Trust–reliance calibration conditions the second-

stage relationship because competent evaluation improves outcomes only when the investor uses the system in proportion to its capability. A residual direct path from design quality to decision quality is estimated rather than assumed away; it may capture reduced search cost, default effects, or unmeasured mechanisms.

Proposed dual-moderated mediation framework

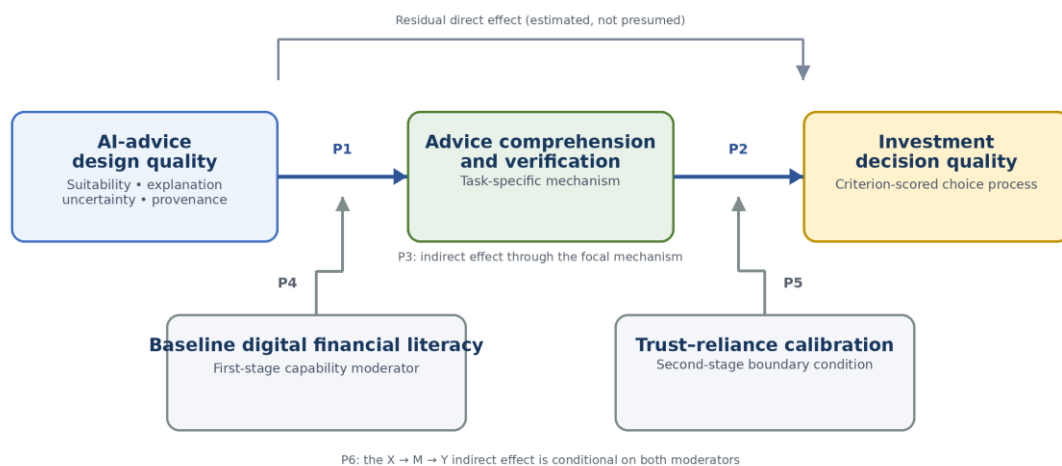


Figure 1. Proposed dual-moderated mediation framework

1. Design quality, comprehension, and decision quality

High-quality advice design directs attention to relevant inputs and makes the recommendation's logic testable. Suitability cues

help the user connect the advice to personal constraints. Uncertainty ranges discourage false precision, while provenance and recency make material claims checkable. These features should improve comprehension and stimulate verification,

especially when the interface requires an active response rather than passive acknowledgement.

Proposition 1 (P1).

Higher AI-advice design quality is positively associated with advice comprehension and verification during the focal investment decision. Comprehension and verification should, in turn, improve the quality of the final choice. A user who can explain the recommendation, identify the relevant uncertainty, and check suitability is better positioned to accept valid advice and reject invalid advice. The outcome is not blind compliance; justified disagreement can be the highest-quality response when the recommendation is poor.

Proposition 2 (P2).

Advice comprehension and verification are positively associated with criterion-scored investment decision quality. Together, the first two propositions specify a mechanism rather than a correlation among favorable attitudes. The indirect effect is expected because design features alter what the user can understand and verify, and those processes alter the choice. The mediator must therefore be measured before the final decision and separately from perceived design quality.

Proposition 3 (P3).

Advice comprehension and verification mediate the positive relationship between AI-advice design quality and investment decision quality.

2. First-stage moderation by baseline DFL

Design cues do not operate in a cognitive vacuum. Investors with stronger baseline DFL can connect a volatility range to risk capacity, understand why fees compound, and recognize when a recommendation conflicts with diversification. The same disclosure may remain decorative for a user who lacks those concepts. Baseline capability should therefore strengthen the design-quality effect on comprehension and verification. This prediction does not imply that low-DFL investors should receive less support; it identifies the need for adaptive explanation, examples, and skill-building.

A competing compensatory hypothesis deserves explicit testing. If the explanation is sufficiently structured and accessible, the marginal benefit of improved design may be greater among low-DFL users. A factorial study can distinguish the capability-amplification and compensation patterns by estimating nonlinear interactions and avoiding a median split.

Proposition 4 (P4). Baseline digital financial literacy positively moderates the relationship between AI-advice design quality and advice comprehension

and verification, such that the first-stage relationship is stronger at higher baseline DFL.

Second-stage boundary condition: trust–reliance calibration

Comprehension does not guarantee appropriate action. A user may understand valid advice but refuse it because of generalized algorithm aversion, or may understand a limitation yet follow the advice because the interface appears authoritative. The relevant boundary condition is therefore not trust level alone but whether reliance responds to demonstrated reliability and disclosed limits. When calibration is high, comprehension can guide selective reliance; when calibration is poor, both overreliance and underreliance weaken the translation of cognition into decision quality.

Proposition 5 (P5). Trust–reliance calibration positively moderates the relationship between advice comprehension and verification and investment decision quality, such that the relationship is stronger when reliance is better calibrated to system capability.

The complete model predicts a conditional indirect effect. Good design should be most valuable when the investor can interpret its cues and when reliance remains sensitive to the system's actual competence. This is a dual-moderated mediation claim, not a quadratic main-effect claim. Researchers may separately test whether global trust has an inverted-U association with decision quality, but that analysis should not be mislabeled as moderation.

Proposition 6 (P6). The indirect effect of AI-advice design quality on investment decision quality through advice comprehension and verification is jointly conditional on baseline DFL and trust–reliance calibration; the indirect effect is strongest when both are higher.

India as a Theoretical Boundary Context

1. Awareness, participation, and capability heterogeneity

India is theoretically important because access, knowledge, and participation do not move together. The SEBI Investor Survey 2025, whose main report was published in 2026, used a listing survey of 91,950 households and reported that 63% of households were aware of at least one securities-market product, but only 9.5% participated—approximately 3.21 crore of 33.72 crore households (SEBI, 2025). Participation was 15% in urban households and 6% in rural households, reaching 23% in the top nine metros. These differences are not background demographics; they predict variation in the capability needed to interpret and verify AI advice. The same survey found that complexity and information gaps were barriers for

74% of non-investors, risk-and-return concerns for 73%, and trust and transparency issues for 51% (SEBI, 2025). These patterns make both stages of the model consequential. Better explanation may reduce information friction, but an interface that merely simplifies action can exploit rather than solve the knowledge gap. Qualitative Indian evidence suggests that robo-advice may help address some investor biases, while also reinforcing the importance of context-specific design and testing (Bhatia et al., 2020). Sampling only experienced metropolitan app users would compress the variance central to the theory and overstate generalizability.

2. Social and digital information channels

Retail advice occurs inside a wider information ecology. SEBI (2025) reports that 59% of surveyed investors relied on friends, family, and colleagues for securities-market information, 56% used financial influencers, and 34% used online investment communities. Ninety-three percent considered finfluencers moderately or highly credible, and 62% reported making at least some investment decisions based on finfluencer recommendations. These figures do not prove that such advice is inaccurate, but they show why source verification, conflicts, and authority cues must be incorporated into AI-advice research. Conversational systems can reproduce, summarize, or amplify material from this ecology without making provenance visible. They can also generate persuasive explanations in multiple languages and at low cost. In India, multilingual and mobile-first design may widen access, but translation quality, financial terminology, screen size, and local examples can change comprehension. Language and urbanicity should therefore be treated as measurement and treatment-effect questions, not merely control variables. IOSCO's work on online imitative trading likewise underscores the conduct risks created when digital engagement, social influence, and frictionless execution converge (IOSCO, 2025).

3. Regulatory perimeter and accountability

The regulatory status of a system depends on what it does and who provides it. Personalized investment advice falls within the perimeter of the SEBI (Investment Advisers) Regulations, 2013, as amended (SEBI, 2025b), while security-specific research and recommendations may engage the Research Analysts framework (SEBI, 2025c). An AI interface does not erase the responsibilities of a regulated entity for suitability, disclosures, records, supervision, and grievance redressal. Nor should a generic system be described as a registered adviser merely because a user requests a portfolio

suggestion. SEBI issued a consultation paper in June 2025 on responsible use of AI and machine learning in Indian securities markets (SEBI, 2025a). As of the literature and regulatory cutoff (15 August 2024), that document is treated as a consultation proposal, not final law. Its emphasis on responsibility, model governance, data, testing, and investor protection is nevertheless directly relevant to research design. Internationally, IOSCO (2021) has emphasized governance and controls for AI/ML used by intermediaries and asset managers. The RBI's discussion of technology and financial services similarly frames innovation alongside operational resilience and customer protection (Das, 2020).

The regulatory implication of the proposed model is narrow but important: disclosure should be evaluated by whether it enables comprehension and verification, and trust should be evaluated by whether it produces appropriate reliance. A compliance notice that users cannot interpret is not equivalent to effective transparency. Equally, a trust-building interface that obscures limitations may increase rather than reduce harm.

Empirical Research Design

1. Multi-stage, repeated-trial experiment

A credible test requires temporal separation and objective variation. Stage 1 should measure baseline DFL, financial experience, risk capacity, numeracy, self-confidence, prior AI use, and general propensity to trust automation. Stage 2 should randomize advice-design features in standardized portfolio scenarios. A practical $2 \times 2 \times 2$ design could vary explanation transparency, uncertainty communication, and provenance/recency disclosure while holding the underlying recommendation constant. A separate reliability manipulation should create valid and invalid advice across repeated trials. Stage 3 should measure comprehension and verification before the participant makes an incentivized or consequence-framed choice. Stage 4 should assess delayed recall, transfer to a new scenario, and, where consent permits, subsequent process behavior. Repeated trials are necessary for calibration. Participants cannot calibrate to a system whose capability is never observable. Reliability can vary by task domain, allowing researchers to test whether users learn appropriate boundaries or form a single global impression. The design should record advice acceptance, modifications, requests for evidence, human escalation, and time spent on verification. A transparent system should not be judged successful merely because acceptance rises; success means greater discrimination between good and bad advice.

2. Measurement and scoring

Baseline DFL should combine objective financial questions, digital-risk knowledge, source and regulatory recognition, and applied tasks. The mediator should include comprehension of the recommendation and uncertainty, detection of a deliberately inconsistent input or stale source, and actual verification actions. Attitudinal trust can be measured with task-specific competence, predictability, process, and purpose items. Reliance is behavioral. Calibration can be modeled as trial-level sensitivity to advice validity or summarized as appropriate acceptance plus appropriate rejection, with overreliance and underreliance reported separately. Decision-quality scoring should be developed before data collection with finance and consumer-protection experts. For a retirement scenario, for example, the rubric may score risk-capacity fit, diversification, liquidity, fee burden, and rejection of leverage. For a short-horizon goal, liquidity and downside protection may receive greater weight. Blinded raters or deterministic scoring rules should be used, and inter-rater reliability should be reported where judgment remains. Realized return can be a later secondary outcome, but it should not replace process validity.

3. Sampling and analysis

The Indian sample should be stratified across metropolitan, smaller urban, and rural locations; gender; age; education; investing experience; and language. Cognitive interviews should precede translation, and measurement invariance should be examined before comparing

groups. Oversampling newer and lower-DFL investors is theoretically justified because the national distribution is not well represented by broker-platform convenience samples. Researchers should report device type and screen constraints because a disclosure may function differently on a low-cost mobile device than on a desktop. A multilevel conditional-process model is appropriate because trials are nested within participants. The first-stage equation predicts comprehension and verification from design quality, baseline DFL, and their interaction. The second-stage equation predicts decision quality from the mediator, calibration, and their interaction, while estimating the residual direct effect of design quality. Conditional indirect effects should be calculated with bootstrap or Bayesian interval estimates at substantively meaningful moderator values. Nonlinear terms may test rival compensation or overtrust hypotheses, but they should be prespecified and distinguished from interaction effects. Causal mediation requires assumptions that should be made explicit. Randomizing design quality identifies its effect, but mediator-outcome confounding remains possible. Sequential randomization, encouragement designs, or carefully timed repeated measures can strengthen inference. Self-report common-method variance should be reduced by objective tasks, logs, independent scoring, and separated measurement occasions. Attrition, attention failures, and language-related differential item functioning should be addressed transparently.

Table 3. Recommended design and validity safeguards

Research target	Recommended test	Key safeguard
First-stage mechanism	Randomize design features; measure comprehension and verification before choice	Separate manipulated quality from perceived quality and establish temporal order
Baseline DFL moderation	Pre-exposure objective and applied DFL battery; continuous interaction	Do not treat a stable trait as a same-wave mediator or dichotomize it
Calibration boundary condition	Repeated valid and invalid advice with reliability feedback	Report appropriate reliance, overreliance, and underreliance separately
Decision quality	Incentive-compatible scenarios with a preregistered criterion rubric	Do not substitute adoption intention, acceptance, or short-term return
Indian heterogeneity	Multilingual, stratified sample with measurement-invariance tests	Avoid metropolitan platform users as the only sampling frame
Responsible data use	Minimized, consent-based logs and de-identified analysis files	Predefine access, retention, withdrawal, and incident procedures

Implications

1. Theoretical implications

The framework changes the dependent variable from acceptance to quality. This matters because the same interface feature can increase adoption while weakening scrutiny. It also locates

stable DFL where it belongs temporally as capability and boundary condition—while using advice comprehension and verification as the immediate mechanism. The model therefore supports stronger causal tests than a single survey in which all variables are favorable perceptions. The

trust distinction provides a second contribution. High trust is not inherently beneficial, and low trust is not inherently irrational. Appropriate reliance depends on task-specific evidence about system competence. By modeling trust–reliance calibration, the framework can explain both informed rejection and selective acceptance. It also prevents conceptual slippage between a moderation hypothesis and a quadratic main effect. Finally, the Indian context generates testable heterogeneity. DFL, language, information source, urbanicity, and investing experience should alter the effectiveness of explanation and verification features. India is thus a setting in which capability amplification, design compensation, and social-source substitution can be compared rather than a geographic label added after theory development.

2. Platform and advisory implications

Platforms should design for comprehension, contestability, and calibrated reliance. Recommended features include visible data dates and sources; plain-language uncertainty; suitability checks tied to horizon and risk capacity; fee and conflict disclosures at the decision point; editable inputs; side-by-side alternatives; friction before high-risk orders; and a documented route to qualified human review. Explanations should help the investor predict when the system may be wrong, not simply make the recommendation more persuasive. Providers should validate the full decision pathway. Predictive performance, user comprehension, verification, reliance errors, and subgroup effects are separate tests. A “human in the loop” label is not evidence of safety unless the human has defined responsibility, competence, information, and escalation authority. Conversely, an automated system should not use a long disclosure to transfer responsibility for defects that the provider can control. DFL support should be embedded at relevant moments without manipulating users toward more activity. Micro-learning can explain volatility, diversification, leverage, fees, and source checks when those concepts become decision-relevant. Progress should be assessed through error detection and transfer to new tasks, not clicks or self-reported confidence.

3. Policy and investor-education implications

Policy should target appropriate reliance rather than maximum digital adoption. Minimum governance expectations for regulated uses of AI should cover accountable ownership, documented purpose, data lineage, validation, change control, cybersecurity, bias assessment, suitability, record keeping, incident reporting, vendor oversight, and complaint handling. Investor-facing communications should state whether content is

generated or verified, whether it is personalized, which entity is responsible, and what material incentives or conflicts exist. Investor education should use multilingual simulations that mirror actual screens and pressures. A module on AI-enabled advice should test whether a learner can identify an unregistered source, challenge an unsuitable risk score, detect stale evidence, compare fees, and respond to a fabricated performance claim. Regulators, exchanges, educational institutions, and intermediaries can evaluate these programs with common scenario banks. The outcome should be safer decisions and better error detection, not course completion alone.

Limitations and Future Research

This is a conceptual paper and does not report original participants, causal estimates, or a meta-analysis. The synthesis is problem-driven rather than exhaustive, and some evidence comes from settings outside India. The category of AI-enabled advice remains heterogeneous; rules-based robo-advisers and generative systems should not be assumed to share reliability, explanation, or regulatory characteristics. The model also simplifies dynamics across time: repeated use can change DFL, trust, risk perception, and provider behavior. Several extensions follow. Longitudinal research should test whether repeated scaffolded verification produces durable DFL gains or merely temporary compliance. Field experiments should compare adaptive explanations across Indian languages and device constraints. Studies should examine when human escalation improves quality and when it adds authority without value. Research on social-source substitution should test whether users treat AI output as independent confirmation when it merely summarizes the same influencer or community content. Finally, welfare analysis should connect decision-process quality to longer-run costs, diversification, and suitability without attributing market noise to the quality of the original choice.

Conclusion

AI-enabled financial advice can broaden access to analysis, but access, adoption, and welfare are not interchangeable. A defensible theory must explain what occurs between an interface and an investment choice. The proposed model identifies advice comprehension and verification as that task-specific mechanism, baseline digital financial literacy as the capability that conditions it, and trust–reliance calibration as the boundary condition that determines whether evaluation becomes appropriate action. For Indian retail investors, these relationships are especially consequential because high awareness coexists with low and uneven

participation, substantial information and trust barriers, and heavy use of social and finfluencer channels. Responsible innovation therefore requires transparent and contestable design, criterion-based evaluation, multilingual capability support, and accountability that remains with the relevant provider. The six propositions and empirical design advanced here offer a coherent basis for testing when AI advice improves investment decision quality—and when a persuasive digital interface merely accelerates a poor decision.

Acknowledgement

The authors would like to express their sincere gratitude to all those who contributed, directly or indirectly, to the development of this conceptual paper, “AI-Enabled Financial Advice and Investment Decision Quality: A Dual-Moderated Mediation Framework for Indian Retail Investors.” The authors gratefully acknowledge the valuable contributions of the scholarly literature on artificial intelligence, robo-advisory, digital financial literacy, information adoption, human–automation trust, and investment decision-making that provided the theoretical foundation for this study.

Financial support and sponsorship

Nil.

Conflicts of interest

The authors declare that there are no conflicts of interest regarding the publication of this paper.

References

1. 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)
2. Belanche, D., Casaló, L. V., & Flavián, C. (2019). Artificial intelligence in FinTech: Understanding robo-advisors adoption among customers. *Industrial Management & Data Systems*, 119(7), 1411–1430. <https://doi.org/10.1108/IMDS-08-2018-0368>
3. Bhatia, A., Chandani, A., & Chhateja, J. (2020). Robo advisory and its potential in addressing the behavioral biases of investors—A qualitative study in Indian context. *Journal of Behavioral and Experimental Finance*, 25, Article 100281. <https://doi.org/10.1016/j.jbef.2020.100281>
4. Buçinca, Z., Malaya, M. B., & Gajos, K. Z. (2021). To trust or to think: Cognitive forcing functions can reduce overreliance on AI in AI-assisted decision-making. *Proceedings of the ACM on Human-Computer Interaction*, 5(CSCW1), Article 188, 1–21. <https://doi.org/10.1145/3449287>
5. Cardillo, G., & Chiappini, H. (2024). Robo-advisors: A systematic literature review. *Finance Research Letters*, 62, Article 105119. <https://doi.org/10.1016/j.frl.2024.105119>
6. Chhillar, N., Arora, S., & Chawla, P. (2024). Measuring digital financial literacy: Scale development and validation. *Thailand and the World Economy*, 42(1), 110–145. <https://so05.tci-thaijo.org/index.php/TER/article/view/270077>
7. Das, S. (2020, February 24). Banking landscape in the 21st century [Speech]. Reserve Bank of India. <https://www.rbi.org.in/commonman/english/Scripts/speeches.aspx?Id=3160>
8. 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>
9. 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>
10. Dietvorst, B. J., Simmons, J. P., & Massey, C. (2018). Overcoming algorithm aversion: People will use imperfect algorithms if they can (even slightly) modify them. *Management Science*, 64(3), 1155–1170. <https://doi.org/10.1287/mnsc.2016.2643>
11. 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>
12. Hoff, K. A., & Bashir, M. (2015). Trust in automation: Integrating empirical evidence on factors that influence trust. *Human Factors*, 57(3), 407–434. <https://doi.org/10.1177/0018720814547570>
13. International Organization of Securities Commissions. (2021). The use of artificial intelligence and machine learning by market intermediaries and asset managers (Final Report). <https://www.iosco.org/library/pubdocs/pdf/IOSCPD684.pdf>
14. International Organization of Securities Commissions. (2025). Online imitative trading practices: Copy trading, mirror trading, social trading (Final Report FR/06/2025). <https://www.iosco.org/library/pubdocs/pdf/IOSCPD793.pdf>
15. Lee, J. D., & See, K. A. (2004). Trust in automation: Designing for appropriate reliance. *Human Factors*, 46(1), 50–80. https://doi.org/10.1518/hfes.46.1.50_30392

16. 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>
17. Lourenço, C. J. S., Dellaert, B. G. C., & Donkers, B. (2020). Whose algorithm says so: The relationships between type of firm, perceptions of trust and expertise, and the acceptance of financial robo-advice. *Journal of Interactive Marketing*, 49, 107–124. <https://doi.org/10.1016/j.intmar.2019.10.003>
18. Lusardi, A., & Mitchell, O. S. (2014). The economic importance of financial literacy: Theory and evidence. *Journal of Economic Literature*, 52(1), 5–44. <https://doi.org/10.1257/jel.52.1.5>
19. Namyslo, N. M., Jung, D., & Sturm, T. (2025). The state of robo-advisory design: A systematic consolidation of design requirements and recommendations. *Electronic Markets*, 35, Article 20. <https://doi.org/10.1007/s12525-025-00762-2>
20. OECD. (2022). OECD/INFE toolkit for measuring financial literacy and financial inclusion 2022. OECD Publishing. <https://doi.org/10.1787/cbc4114f-en>
21. OECD. (2023). OECD/INFE 2023 international survey of adult financial literacy. OECD Business and Finance Policy Papers, No. 39. OECD Publishing. <https://doi.org/10.1787/56003a32-en>
22. OECD. (2025). Artificial intelligence and personal finance. OECD Artificial Intelligence Papers, No. 62. OECD Publishing. <https://doi.org/10.1787/2858fdf4-en>
23. Oehler, A., & Horn, M. (2024). Does ChatGPT provide better advice than robo-advisors? *Finance Research Letters*, 60, Article 104898. <https://doi.org/10.1016/j.frl.2023.104898>
24. Piehlmaier, D. M. (2022). Overconfidence and the adoption of robo-advice: Why overconfident investors drive the expansion of automated financial advice. *Financial Innovation*, 8, Article 14. <https://doi.org/10.1186/s40854-021-00324-3>
25. Ravikumar, T., Suresha, B., Prakash, N., Vazirani, K., & Krishna, T. A. (2022). Digital financial literacy among adults in India: Measurement and validation. *Cogent Economics & Finance*, 10(1), Article 2132631. <https://doi.org/10.1080/23322039.2022.2132631>
26. Securities and Exchange Board of India. (2025a). Consultation paper on guidelines for responsible usage of AI/ML in Indian securities markets. https://www.sebi.gov.in/reports-and-statistics/reports/jun-2025/consultation-paper-on-guidelines-for-responsible-usage-of-ai-ml-in-indian-securities-markets_94687.html
27. Securities and Exchange Board of India. (2025b, November 25). Securities and Exchange Board of India (Investment Advisers) Regulations, 2013 (as amended). https://www.sebi.gov.in/legal/regulations/nov-2025/securities-and-exchange-board-of-india-investment-advisers-regulations-2013-and-securities-last-amended-on-november-25-2025_98246.html
28. Securities and Exchange Board of India. (2025c, November 25). Securities and Exchange Board of India (Research Analysts) Regulations, 2014 (as amended). https://www.sebi.gov.in/legal/regulations/nov-2025/securities-and-exchange-board-of-india-research-analysts-regulations-2014-last-amended-on-november-25-2025_98248.html
29. Securities and Exchange Board of India. (2025). Investor Survey 2025: Main report. https://www.sebi.gov.in/sebi_data/commndoc/s/jan-2026/Investor%20Survey%202025%20Main%20Report.pdf
30. Snyder, H. (2019). Literature review as a research methodology: An overview and guidelines. *Journal of Business Research*, 104, 333–339. <https://doi.org/10.1016/j.jbusres.2019.07.039>
31. Sussman, S. W., & Siegal, W. S. (2003). Informational influence in organizations: An integrated approach to knowledge adoption. *Information Systems Research*, 14(1), 47–65. <https://doi.org/10.1287/isre.14.1.47.14767>
32. Wagner, F. (2024). Determinants of conventional and digital investment advisory decisions: A systematic literature review. *Financial Innovation*, 10, Article 18. <https://doi.org/10.1186/s40854-023-00538-7>
33. Zhu, H., Sallnäs Pysander, E.-L., & Söderberg, I.-L. (2023). Not transparent and incomprehensible: A qualitative user study of an AI-empowered financial advisory system. *Data and Information Management*, 7(3), Article 100041. <https://doi.org/10.1016/j.dim.2023.100041>
34. Zhu, H., Vigren, O., & Söderberg, I.-L. (2024). Implementing artificial intelligence empowered financial advisory services: A literature review and critical research agenda. *Journal of Business Research*, 174, Article 114494. <https://doi.org/10.1016/j.jbusres.2023.114494>