

Advice Without an Advisor: Generative AI, Epistemic Deference, and the Calibration of Trust in Household Economic Decisions.

Vishal Kaushal¹, Akanksha Negi²

¹Senior Research Fellow, Department of Commerce, Himachal Pradesh University, Shimla

²Senior Research Fellow, Department of Commerce, Himachal Pradesh University, Shimla

Cite this paper as: Vishal Kaushal, Akanksha Negi (2026) Advice Without an Advisor: Generative AI, Epistemic Deference, and the Calibration of Trust in Household Economic Decisions... Journal of International Marketing and Marketing Research, 4(2) 7-12

ABSTRACT

Households increasingly consult general-purpose generative artificial intelligence (AI) systems for economic decisions, from budgeting and debt to saving and investing, inserting a fluent, confident, and unaccountable interlocutor into reasoning that was previously done alone, with a human professional, or not at all. This article asks how people calibrate trust in such advice, and argues that generative advisors break the cues on which the existing science of advice-taking was built. That science, spanning algorithm aversion and algorithm appreciation, was developed largely with narrow numeric-estimation tasks in which an advisor was labelled as a person or an algorithm and its output was brief and impersonal. Generative advisors differ on three dimensions that matter for calibration. They personalize through dialogue, which should neutralize the uniqueness neglect that has driven resistance to earlier medical and consumer AI. They express fluency and confidence that are decoupled from accuracy, so the surface signals people use to gauge competence no longer track it. And they carry none of the fiduciary, relational, or reputational accountability that a human financial advisor bears, yet nothing in the interaction makes that absence salient. The article integrates the advice-taking, algorithm-reliance, trust-in-automation, and financial-advice literatures into a conceptual model of miscalibrated deference, expressed as six propositions, and locates the problem in a Global South context where the scarcity of human financial advice means generative systems fill a vacuum rather than compete with an incumbent. Because the mechanisms are novel and partly interior, the article specifies an exploratory sequential mixed-methods design: a qualitative phase combining interviews with think-aloud sessions on participants' own financial questions, followed by a vignette experiment using the judge–advisor paradigm to quantify deference across manipulated conditions. The article reframes AI financial advice as a problem of trust calibration rather than of output quality, and specifies how to study it...

Keywords: *generative AI, financial advice, algorithm aversion, trust calibration, advice taking, epistemic deference.....*

Introduction

A person deciding whether to prepay a home loan, which fund to choose, or how to clear a credit-card balance now has a new place to ask. Instead of consulting a bank, a relative, or no one, they can open a chat window and put the question to a general-purpose AI system that answers immediately, in fluent prose, with the calm assurance of an expert and the patience of a machine. Studies of these systems have shown that recent models perform well on financial-literacy tests and can produce plausible, sometimes usefully specific personal-finance guidance (Niszczota & Abbas, 2023). The natural research question that follows has been whether the advice is good. This article argues that a prior and more consequential question has been underexamined: whether people can tell, and how much they defer regardless.

Deference is the crux. Advice only matters to a decision insofar as the decision-maker takes it, and decades of research show that advice-taking is neither automatic nor well-calibrated. People discount advice egocentrically, staying closer to their own initial view than accuracy would warrant (Bonaccio & Dalal, 2006; Yaniv & Kleinberger, 2000), yet under other conditions they over-rely, following automated aids past the point of good judgment (Parasuraman & Riley, 1997). Which way the error runs depends on how the advisor is perceived, and here generative AI presents something the existing literature did not study. It is not the labelled statistical algorithm of the algorithm-aversion experiments, nor the impersonal recommender of a robo-advisor. It is a conversational agent that appears to understand the specific person it is talking to, speaks with unwavering fluency, and answers for nothing.

These three properties, conversational personalization, fluency decoupled from accuracy, and the absence of accountability, together dismantle the cues people have historically used to calibrate how much to trust advice. The result is not simply more trust or less, but miscalibrated trust: deference that no longer tracks what it should. The article develops this claim by integrating four literatures that have largely run in parallel, derives a set of propositions, and specifies a research design suited to investigating them. It also attends to a setting the dominant literature has neglected. In much of India, professional.

....:

financial advice is scarce, and a generative advisor does not compete with a trusted human incumbent but substitutes for the absence of one, which raises both the likely level of deference and the stakes of error

2. From Robo-Advisors to Conversational General-Purpose Advisors

Automated financial advice is not new, but its earlier forms were narrow. Robo-advisors allocate portfolios from a questionnaire; recommender systems suggest products; rule-based tools answer scoped questions. Their outputs are constrained, their domains explicit, and their limits legible to the user, who knows they are interacting with a specialized instrument. Generative systems differ in kind. They accept any question in natural language, respond across every domain of household finance at once, and do so in the register of a knowledgeable human rather than a form or a dashboard. Recent evaluations establish that the underlying capability is real and improving: contemporary models pass financial-literacy assessments and generate advice that experts often judge reasonable, if uneven and occasionally wrong in ways a novice would not catch (Niszczota & Abbas, 2023). Working papers exploring model-generated retirement and life-cycle guidance likewise find the outputs frequently sensible while flagging systematic gaps (Lo & Ross, 2024).

The behavioural question does not follow from the capability question. Even excellent advice fails if it is ignored, and even poor advice harms if it is followed. What governs the gap is the psychology of reliance, and the conversational, general-purpose form of generative advice changes the inputs to that psychology in ways the robo-advisor never did. To see how, the relevant behavioural science must be assembled, because no single strand of it was built for this object.

3. Theoretical Framework

3.1 Advice taking and the weighting of advice

The judge–advisor tradition treats decision-making as a two-stage process: a person forms an initial judgment, receives advice, and revises. The degree of revision, operationalized as the weight of advice, is the central measure, and the robust finding is that people underweight advice relative to what accuracy would dictate, remaining anchored on their own first opinion (Bonaccio & Dalal, 2006; Yaniv & Kleinberger, 2000). This egocentric discounting is not uniform. Advice is weighted more when the advisor is seen as expert, when the task is hard, and when the decision-maker is uncertain, and less when the advisor’s reputation has been damaged. The framework matters here because it gives a precise vocabulary and a measurable quantity for deference, and because it locates the action in perceptions of the advisor rather than in the advice itself.

3.2 Algorithm aversion and Algorithm appreciation

When the advisor is an algorithm, the picture is contested in an instructive way. One line of work finds algorithm aversion: people abandon an algorithmic forecaster more readily than a human one after seeing it err, even when it outperforms the human, because they lose confidence in the machine faster (Dietvorst et al., 2015). A complementary line finds the opposite under different conditions, an algorithm appreciation in which people weight identical advice more heavily when told it came from an algorithm than from a person (Logg et al., 2019). The reconciliation is that reliance is domain-specific and cue-dependent: aversion intensifies after visible error and in tasks seen as human, appreciation appears in tasks seen as objective, and small design features, such as letting people adjust the algorithm’s output, can flip the response (Dietvorst et al., 2018; Prah & Van Swol, 2017; Singh et al., 2022a). The lesson for generative advisors is that there is no fixed disposition to be recovered; how much people defer depends on how the system presents itself, and generative systems present themselves very differently from the terse forecasters of these studies.

3.3 Uniqueness neglect and its Possible undoing

A specific mechanism from consumer research is central to my argument. Longoni et al. (2019) showed that consumers resist medical AI because of uniqueness neglect, the belief that an automated provider, however accurate on average, cannot accommodate what is distinctive about their own case. Resistance was strongest among those who saw themselves as more unique, and it dissolved when the AI’s care was framed as personalized. A parallel result in marketing finds that receptivity to machine advice depends on whether the task is seen as utilitarian or as bound up with subjective, identity-laden judgment, the so-called word-of-machine effect (Longoni & Cian, 2022; Singh et al., 2022b). Uniqueness neglect has plausibly restrained reliance on earlier financial automation, which asks standardized questions and returns standardized allocations. The pivotal observation is that a conversational generative advisor is built to undo exactly this restraint. It asks about the person’s situation, incorporates their particulars into its reply, and addresses them by the details they supplied. If uniqueness neglect is what held reliance in check, conversational personalization is designed, whether or not by intent, to remove the brake.

3.4 Trust calibration and the Accountability void

Finally, the automation-trust literature reframes the goal. The aim is not maximal trust but appropriate reliance, trust

calibrated to the system's actual capabilities, with miscalibration taking two forms: disuse of good automation and misuse of bad (Lee & See, 2004; Parasuraman & Riley, 1997). Calibration requires cues that track capability. Two features of generative advice corrupt those cues. First, fluency and expressed confidence are decoupled from accuracy: the system writes with the same assurance whether it is right or wrong, so the verbal markers people use to infer competence in human advisors no longer carry information. Second, and distinctively, the advisor bears no accountability. Research on human financial advice describes advisors partly as money doctors whose value lies as much in the trust relationship as in the information conveyed, a relationship backed by reputation, regulation, and the advisor's stake in the client (Gennaioli et al., 2015; Singh et al., 2018). A generative advisor has none of this: no fiduciary duty, no reputation to lose with the specific user, no exposure to the consequences of its counsel. Nothing in the fluent, personalized exchange makes that absence felt, and there is little reason to expect users to supply the discount on their own.

These strands converge on a thesis. The cues that made advice-taking even roughly calibratable, an advisor whose competence showed in its manner, whose standardization signalled its limits, and whose accountability underwrote its counsel, are precisely the cues that generative advice removes or inverts. The problem is therefore not reducible to output quality. It is a problem of calibration under corrupted cues.

4. Conceptual Model of Miscalibrated Deference

The framework advanced here holds that deference to generative economic advice will systematically depart from warranted reliance, and that this departure has identifiable sources and an identifiable distribution. The propositions below state the framework. They are directional statements advanced to guide empirical investigation, and each is open to disconfirmation.

Proposition 1. Conversational personalization raises deference to generative advice above the level accorded to functionally equivalent but impersonal automated advice, because dialogic tailoring neutralizes the uniqueness neglect that restrains reliance on standardized systems.

Proposition 2. The fluency and expressed confidence of generative advice raise deference independent of its accuracy, so that trust is driven by surface features that do not track competence.

Proposition 3. Users under-weight the advisor's lack of accountability: deference does not adjust downward for the absence of fiduciary duty, reputation, and stake in the outcome, because nothing in the interaction renders that absence salient.

Proposition 4. Deference is domain-moderated: it is higher when the decision is framed as technical or utilitarian and lower when it is experienced as identity-laden or moral, extending the domain-specificity found in prior algorithm-reliance and word-of-machine research to conversational advisors.

Proposition 5. Where access to human financial advice is scarce, generative advisors substitute for absent advisors rather than competing with present ones; deference is therefore higher and less checked by an available human benchmark, making the calibration problem more consequential in such settings than in well-advised populations.

Proposition 6. Miscalibration is distributed unequally: users with lower financial and AI literacy exhibit larger gaps between their deference and the reliance the advice warrants, concentrating the resulting harms among those least able to absorb them.

The framework's contribution is to shift the object of concern. If the difficulty with AI financial advice were only that models sometimes err, the remedy would be better models. The account developed here implies that even accurate systems can produce harm through the reliance they invite, so that the locus of both risk and intervention is the calibration of trust rather than the quality of output alone. The framework also makes a distributive prediction, in Proposition 6, that turns a general question about human-AI interaction into a question about who is exposed, which is where the Global South framing becomes analytically consequential rather than merely contextual.

5. An Exploratory Sequential Mixed-Methods Design

The mechanisms in the model are partly interior, how deference is formed in the moment of consulting the system, and partly measurable, how much people revise their decisions under manipulated conditions. This combination calls for an exploratory sequential mixed-methods design, in which a qualitative phase surfaces and specifies the constructs and a subsequent quantitative phase tests them experimentally (Creswell & Plano Clark, 2018). The qualitative phase leads because the object is new and the relevant constructs should be grounded in how people actually reason with these systems rather than imposed a priori.

5.1 Phase 1: Qualitative Elicitation

The first phase combines semi-structured interviews with think-aloud sessions. Participants would bring a real, personally relevant financial question and work through it with a generative system while verbalizing their reasoning, following established protocol-analysis procedure, in which concurrent verbalization is treated as a window onto the reasoning in

progress rather than as retrospective justification (Ericsson & Simon, 1993; Kumar et al., 2026a; Kaushal et al., 2026b; Singh et al., 2024). The value of the think-aloud is that it captures the moment of calibration: where the participant accepts a claim, where they hesitate, what makes them check or fail to check, and how the system's personalization and confidence enter their acceptance. Interviews before and after would situate the session in the participant's financial circumstances, prior sources of advice, and access to human professionals, the last being essential for the Global South argument.

Sampling would be purposive and designed for contrast (Palinkas et al., 2015): across financial literacy, across prior AI use, and crucially across access to human financial advice, so that the substitution dynamic in Proposition 5 can be observed rather than assumed. Analysis would use reflexive thematic analysis (Braun & Clarke, 2006), with the researcher's interpretive role treated as part of the method rather than a contaminant, and would be evaluated against interpretive trustworthiness criteria (Lincoln & Guba, 1985). The phase has two deliverables: a grounded account of how deference forms, and an operational apparatus, constructs, framings, and stimulus materials, for the experiment that follows.

5.2 Phase 2: A Judge–Advisor Experiment

The second phase quantifies deference using the judge–advisor paradigm and its weight-of-advice measure (Bonaccio & Dalal, 2006; Kumar et al., 2026b; Kaushal et al., 2026a; Singh et al., 2025; Yaniv & Kleinberger, 2000). Participants would make an initial judgment on a realistic financial scenario, receive advice, and revise; the shift toward the advice indexes deference. The design would manipulate the features the model implicates. Source would vary across a human advisor, a generative AI advisor, and an impersonal automated tool, holding the content of the advice constant, to isolate the effect of the advisor's identity. Personalization would vary between a generic reply and a conversationally tailored one, to test Proposition 1. Expressed confidence and fluency would vary independently of the advice's correctness, to test Proposition 2. Accountability cues would be manipulated by making the advisor's stake, or its absence, salient, to test Proposition 3. Domain would vary between utilitarian and identity-laden framings of otherwise matched decisions, to test Proposition 4. Financial and AI literacy would be measured as individual differences, to test Proposition 6, and access to human advice would be measured and, where possible, sampled for, to test Proposition 5.

Because this phase is confirmatory in intent, it should be pre-registered, with hypotheses, manipulations, exclusion rules, and analysis plan specified in advance, and powered on effect-size estimates drawn from the advice-taking literature rather than assumed. Leading with the qualitative phase grounds the manipulations and measures in how deference actually forms, reducing the risk of a well-powered test of a misspecified construct.

5.3 Ethical Considerations

The study touches real financial decisions, which imposes obligations. No phase should deliver actual investment or debt advice to participants or leave them acting on stimulus material; scenarios used for measurement must be clearly bounded as research tasks and debriefed, and any personal question worked through in Phase 1 must be accompanied by a clear statement that the session is not financial advice and by signposting to legitimate sources. Manipulating confidence and personalization to study over-reliance carries a duty not to instill it; debriefing must correct any miscalibration the procedure may have encouraged. Given Proposition 6, recruitment among lower-literacy participants requires particular care with comprehension, consent, and the avoidance of harm.

6. Discussion and Implications

The article contributes, first, a reframing: it moves the study of AI financial advice from the quality of the output to the calibration of the user's trust, arguing that harm can flow from accurate systems through the reliance they invite. Second, it integrates four literatures, advice-taking, algorithm reliance, automation trust, and financial advice, that have addressed pieces of this problem in isolation, and shows that generative advisors sit precisely at their intersection while fitting none of them cleanly. Third, it identifies specific mechanisms, conversational personalization undoing uniqueness neglect, fluency decoupled from accuracy, and an unmarked accountability void, that are novel to the generative case and absent from the paradigms that preceded it. Fourth, once executed, the design would yield both a grounded account of how deference forms in situ and experimental estimates of what drives it, in a population where generative advice substitutes for scarce human advice rather than competing with it, extending a literature built largely on well-advised Northern samples.

6.1 Limitations and Boundary Conditions

Several limits bound the argument. The framework concerns household economic decisions of moderate stakes and complexity; very high-stakes decisions, where people may seek human advice regardless, and trivial ones, where calibration hardly matters, fall outside it. Generative systems evolve rapidly, and both their accuracy and their conversational manner shift across model versions, so any experimental estimate is time-bound; the conceptual contribution is intended to outlast particular models, but the empirical parameters will not. The propositions are directional and refutable, and Proposition 1 in particular could fail if a residual scepticism toward machines survives personalization, or reverse if conversational tailoring instead heightens salience of the system's non-humanness; the design is built to detect either outcome. Finally, the think-

aloud method changes the very reasoning it observes, and the experiment trades realism for control; the mixed-methods structure is intended to let each phase compensate for the other's weakness, rather than to present either as sufficient on its own.

7. Conclusion

A confident, fluent, endlessly patient advisor that appears to understand a person's particular situation and answers for nothing is a new kind of object to reason with, and the natural first question, whether its advice is any good, is not the question that most determines whether people are helped or harmed. What determines that is how much they defer, and generative systems arrive having quietly removed or inverted the cues that once made deference calibratable. Personalization dissolves the sense that a machine cannot grasp one's particular case; fluency substitutes for the harder-won signals of competence; and the absence of accountability goes unmarked in an exchange that feels like counsel. This article has argued that the resulting miscalibration is the central concern, has specified the mechanisms that produce it, and has set out a way to study it that respects both the interior character of trust and the demand for measurable evidence. In settings where a generative advisor is not one option among many but the only advisor a household can reach, calibrating that trust correctly is not an academic nicety. It bears on whether a powerful new tool narrows or widens the distance between those who can afford good advice and those who cannot...

REFERENCES

- [1] Bonaccio, S., & Dalal, R. S. (2006). Advice taking and decision-making: An integrative literature review, and implications for the organizational sciences. *Organizational Behavior and Human Decision Processes*, 101(2), 127–151. <https://doi.org/10.1016/j.obhdp.2006.07.001>
- [2] Braun, V., & Clarke, V. (2006). Using thematic analysis in psychology. *Qualitative Research in Psychology*, 3(2), 77–101. <https://doi.org/10.1191/1478088706qp063oa>
- [3] Creswell, J. W., & Plano Clark, V. L. (2018). *Designing and conducting mixed methods research* (3rd ed.). Sage.
- [4] 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>
- [5] 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>
- [6] Ericsson, K. A., & Simon, H. A. (1993). *Protocol analysis: Verbal reports as data* (Rev. ed.). MIT Press.
- [7] Gennaioli, N., Shleifer, A., & Vishny, R. (2015). Money doctors. *The Journal of Finance*, 70(1), 91–114. <https://doi.org/10.1111/jofi.12188>
- [8] Kaushal, V., Chauhan, M., & Pathania, K. S. (2026a). Financial inclusion and economic growth: Evidence from emerging economies. <https://doi.org/10.64751/rb1wdm74>
- [9] Kaushal, V., Singh, R., & Pathania, K. S. (2026b). Mapping financial inclusion and economic empowerment research: A bibliometric analysis and PRISMA-based systematic review.
- [10] Kumar, S., Sharma, D., Kumar, S., et al. (2026a). Analyzing stock returns and volatility trends in emerging markets: Insights from India and China. *IUP Journal of Accounting Research & Audit Practices*, 25(3), 42–59.
- [11] Kumar, S., Thalyari, N., & Kaushal, V. (2026b). Stock market dynamics in BRICS and advanced nations: A comparative quantitative analysis. *IUP Journal of Accounting Research & Audit Practices*, 25(3), 106–122.
- [12] 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
- [13] Lincoln, Y. S., & Guba, E. G. (1985). *Naturalistic inquiry*. Sage.
- [14] Lo, A. W., & Ross, J. (2024). Can ChatGPT plan your retirement? Generative AI and financial advice [Working paper]. Massachusetts Institute of Technology.
- [15] 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>
- [16] Longoni, C., & Cian, L. (2022). Artificial intelligence in utilitarian vs. hedonic contexts: The “word-of-machine” effect. *Journal of Marketing*, 86(1), 91–108. <https://doi.org/10.1177/0022242920957347>

-
- [17]Longoni, C., Bonezzi, A., &Morewedge, C. K. (2019). Resistance to medical artificial intelligence. *Journal of Consumer Research*, 46(4), 629–650. <https://doi.org/10.1093/jcr/ucz013>
- [18]Niszczoła, P., & Abbas, S. (2023). GPT has become financially literate: Insights from financial literacy tests of GPT and a preliminary test of how people use it as a source of advice. *Finance Research Letters*, 58, 104333. <https://doi.org/10.1016/j.frl.2023.104333>
- [19]Palinkas, L. A., Horwitz, S. M., Green, C. A., Wisdom, J. P., Duan, N., &Hoagwood, K. (2015). Purposeful sampling for qualitative data collection and analysis in mixed method implementation research. *Administration and Policy in Mental Health and Mental Health Services Research*, 42(5), 533–544. <https://doi.org/10.1007/s10488-013-0528-y>
- [20]Parasuraman, R., & Riley, V. (1997). Humans and automation: Use, misuse, disuse, abuse. *Human Factors*, 39(2), 230–253. <https://doi.org/10.1518/001872097778543886>
- [21]Prahl, A., & Van Swol, L. (2017). Understanding algorithm aversion: When is advice from automation discounted? *Journal of Forecasting*, 36(6), 691–702. <https://doi.org/10.1002/for.2464>
- [22]Singh, R. K., Kumar, A., & Kumari, J. (2022a). An empirical application of gravity model theory to Indo-BIMSTEC business relations. *Journal of Polity and Society*, 14(1), 119–137. <https://journalspoliticalscience.com/index.php/i/article/view/6>
- [23]Singh, R. K., Kumar, A., Kumari, J., & Singh, Y. (2022b). Empirical study of India's trade and investment relations with ASEAN countries: Evidence from the gravity model theory. *Arthshastra Indian Journal of Economics & Research*, 11(3), 8–28. <https://doi.org/10.17010/aijer/2022/v11i3/172680>
- [24]Singh, R. K., Singh, B., & Chauhan, D. (2018). India's trade relations with SAARC countries in the pre and post-SAFTA periods: An evaluation of trade flows by using multiple regression model. *Arthshastra Indian Journal of Economics & Research*, 7(6), 43–59. <https://doi.org/10.17010/aijer/2018/v7i6/141433>
- [25]Singh, R. K., Kumar, A., Kumari, J., & Singh, Y. (2024). Navigating determinants of India business relations with GCC countries: Evidence from the gravity model theory. *Prabandhan: Indian Journal of Management*, 15(6), 8–27. <https://doi.org/10.101710/pjom/v15i6/170021>
- [26]Singh, R. K., Naresh, K., Singh, Y., Kumar, A., & Kumar, S. (2025). Moderating role of stock return volatility in corporate governance-firm value structural linkages: Insights from a PLS-SEM with IPMA approach. *Indian Journal of Finance*, 19(1), 8–32. <https://doi.org/10.101710/ijf/v19i1/174693>
- [27]Yaniv, I., & Kleinberger, E. (2000). Advice taking in decision making: Egocentric discounting and reputation formation. *Organizational Behavior and Human Decision Processes*, 83(2), 260–281. <https://doi.org/10.1006/obhd.2000.2909>
- .
- .