

Artificial Confidence: A Secondary Data Analysis of AI-Labeled Advice and Cognitive Effects Across Age Groups

Kanishk Mahasuar

Ballarat Clarendon College, 1425 Sturt St, Newington, VIC, Australia 3350; kanishk.mahasuar@gmail.com

ABSTRACT: Artificial intelligence (AI) is being rapidly adopted in decision-making in critical sectors such as health, education, and public services. However, there is a paucity of research assessing the differential impact of AI- or human-labeled advice on human cognition and resulting behavior. We attempted to understand AI's impact on human decision-making through a secondary data analysis of the Human-AI Interactions Dataset, consisting of approximately 34,000 decision-making records. We used two outcome measures: confidence drift (the alteration in self-reported confidence from pre- to post-advice) and decision change (the user's revision of their answer). Results revealed that AI-labeled advice significantly increased confidence drift relative to human advice ($p < .00001$). However, the increase in confidence did not correspond to a higher likelihood of decision change. Age also emerged as a significant negative predictor of confidence drift and decision change. However, the interaction between age and advice source was not statistically significant. The study results suggest that while AI advice may increase users' perceived decision certainty, it does not change actual behavior. Hence, to improve human-AI collaboration, future AI systems should focus on trust calibration, age sensitivity, and the promotion of critical thinking without inducing overconfidence.

KEYWORDS: Behavioral and Social Sciences, Cognitive Psychology, Human-AI Interaction, Decision Making, Trust Calibration.

■ Introduction

The adoption of artificial intelligence (AI) in decision-making processes has accelerated across our society. AI-based decision-making is deployed in healthcare, education, and governance. Arguably, AI systems deliver efficiency and convenience. But they also raise concerns about the influence of AI-generated advice on human responses. One of the important factors determining this response is trust. AI systems develop through a combination of dispositional tendencies, situational context, and learned experience over time, reflecting both cognitive evaluations (e.g., perceived reliability and competence) and affective responses.¹ Trusting AI dictates how individuals rely on AI recommendations, often irrespective of its accuracy.^{2,3} Trust can evolve and may become miscalibrated, leading to either over-reliance on incorrect AI or rejection of correct advice.⁴

Many recent studies have examined cognitive challenges for AI users in decision-making. Srinivasan and Thomason⁵ proposed adaptive systems to recalibrate trust in AI to improve human-AI collaboration. Alon-Barkat and Busuioc⁶ emphasized the importance of users' pre-existing beliefs that bias them to adhere to AI advice. One of the classic works in the field by Yaniv and Kleinberger⁷ demonstrates that individuals incorporate external advice only partially, often ignoring even accurate recommendations in favor of their own initial judgment. This phenomenon is known as egocentric discounting. Persistent cognitive tension between personal judgment and external input from AI should be examined carefully to make human-AI collaboration more trustworthy. However, current literature largely focuses on behavioral outcomes rather than internal psychological effects.⁸ This study addresses that gap

by examining whether AI-labeled advice increases users' confidence more than identical human-labeled advice and whether this effect varies across age groups. Using over 34,000 decision-making records from the Human-AI Interactions Dataset, this paper contributes to understanding how AI affects user cognition by focusing on confidence drift as a key outcome measure.⁹ This study specifically compares AI-labelled and human-labelled advice and does not examine conditions of complete advice rejection, which remain an important area for future research.

Several theoretical models have been developed to explain the variability of trust in AI and its impact on user engagement. Hoff and Bashir¹⁰ proposed a three-layered model of trust variability. It is divided into dispositional (enduring tendency), situational (context-dependent), and learned (based on experience). They emphasize that facilitating appropriate trust is critical to improving safety and productivity in human-automation teams. Arguably, both over-trust (misuse) and under-trust (disuse) can lead to adverse events. Jacovi *et al.*³ introduced "contractual trust," where users evaluate AI systems not only by outcome accuracy but by perceived consistency with expected behavior. They also mention the concept of "warranted trust," which is ethically desirable as compared to "unwarranted trust," which can lead to misuse or disuse.

A significant increase in the use of automation has not always led to comparable improvements in efficiency. Hence, researchers have analyzed the importance of trust calibration in human-AI systems. Dzindolet *et al.*² found that initial positivity bias toward unfamiliar AI systems can erode rapidly. Particularly after encountering errors, unless the AI system provides comprehensible explanations, trust can rapidly di-

minish. While Dzindolet *et al.*² demonstrate that errors reduce trust in automated systems, the influence of error severity is not explicitly examined and may also play an important role, with more serious or high-stakes errors likely to result in greater reductions in trust than minor inaccuracies. On the contrary, knowing why the AI might err was vital, as it made people trust and rely on it more. Parasuraman and Riley⁴ identified four modes of automation interaction: use, misuse, disuse, and abuse. They argued that inappropriate trust accounts for serious automation-related failures. Vodrahalli *et al.*⁹ conducted a large-scale study using the Human-AI Interactions Dataset. Their findings suggested that participants accepted AI-labeled advice 5–6% more often than identical human-labeled advice, with small task-specific variation. Their work focused mostly on advice acceptance rather than confidence modulation or age-based effects.

Beyond system performance, individual differences like demographic factors can also shape how users engage with automation. Czaja *et al.*¹¹ studied the adoption of technology in the adult population (ages 18–91). They found that older adults are significantly less likely to adopt new technologies. Reduced self-efficacy, lower fluid intelligence, and computer anxiety were the factors influencing older adults. Another study by Pak *et al.*¹² had similar findings, highlighting the significance of age. Younger students effectively calibrated their trust with automation reliability, while older adults and cadets exhibited higher overall trust regardless of reliability. These findings emphasise the need for age-aware AI system design.

Taken together, these findings show that trust formation is a complex process, influenced by both the system design and human cognition. Researchers have explored cognitive interventions to correct such overreliance. Bućinca *et al.*¹³ demonstrated that "cognitive forcing interventions" significantly reduced overreliance on AI. The participants tend to make significantly more correct decisions on incorrect AI predictions. They suggested that human cognitive motivation can moderate the effectiveness of explainable AI solutions. Srinivasan and Thomason⁵ found that both excessively low and high trust levels lead to inappropriate AI reliance. Doctors accepted 26% of AI misdiagnoses when trust was high versus only 8% when trust was lower. They also found that providing supporting and counter-explanations resulted in a 38% reduction in inappropriate reliance and a 20% improvement in decision accuracy. Existing literature establishes the importance of trust and advice-taking in human-AI interactions. The development of trust may also depend on an individual's prior knowledge or expertise in the domain, as users with limited understanding may be more likely to either over-rely on or under-trust AI advice due to uncertainty. But most of them didn't examine internal psychological effects sufficiently. This study stresses confidence drift (change in the participant's confidence) as a key outcome measure, giving a more nuanced perspective on AI's impact on user cognition in decision support systems. This study tests the following hypotheses: H1: AI-labelled advice will produce greater confidence drift than human-labelled advice. H2: AI-labelled advice will increase the probability of

decision change. H3: Age moderates the effect of advice source on both confidence drift and decision change.

■ Methods

We conducted a secondary data analysis of a publicly available Human-AI Interactions Dataset developed by Vodrahalli *et al.*⁹ This study examines how AI-labeled advice affects confidence drift and whether age moderates this cognitive response, both of which remain underexplored.

Dataset and Participants:

The dataset comprised 35,670 decision-making trials contributed by 1,141 adults. The participants were recruited through Amazon Mechanical Turk. They were asked to complete a series of cognitive (decision-making) tasks. Tasks were spread across four domains: art authentication, city prediction, census prediction, and sarcasm detection. The participants had one initial response, and then they viewed the advice that was randomly labeled as either "AI-generated" or "human-provided". Participants also provided a confidence rating for their initial response. They were then shown advice that was randomly labelled as either 'AI-generated' or 'human-provided', although the underlying advice content was identical across conditions. After viewing the advice, they submitted a final response. A total of 34,783 valid interactions were used for analysis after excluding incomplete responses. This design isolates the effect of advice labelling independent of content.

Variables:

The study used two key outcome measures, namely confidence drift and decision change. Confidence drift was defined as the numerical difference between post-advice and pre-advice confidence ratings (post – pre), where positive values indicate an increase in confidence after receiving advice. Decision change was defined as a binary variable (0 = no change, 1 = change), indicating whether participants altered their initial response after receiving advice. Two key predictors used were advice source (labeled as AI or human) and age. Age was measured both as a continuous and categorical variable grouped into six ranges (25 years or younger, 25–34, 35–44, 45–54, 55–64, and 65 years or older).

Statistical Analysis:

We used Datatab (www.datatab.net), an online platform for statistical analysis. Sample characteristics and outcome distributions were summarized with descriptive statistics. The difference in mean confidence drift between AI and human advice groups was measured by an independent samples t-test. A two-way ANOVA was used to examine the interaction effect between advice source and age group on confidence drift. Multiple linear regression was used to predict confidence drift as a function of age, advice type, and their interaction. Finally, a logistic regression model assessed the likelihood of decision change based on age and advice source.

Ethical Considerations:

Dataset by Vodrahalli *et al.*⁹ is licensed under the Creative Commons Attribution 4.0 License. No additional ethics approval was required for this secondary analysis since the dataset is anonymized and previously collected for research purposes.

Results

This study analyzed 34,783 interaction records from the Human-AI Interactions Dataset to examine the influence of AI versus human advice on user confidence and decision changes.

Confidence Drift by Advice Source:

An independent samples t-test was conducted to compare confidence drift between participants who received advice labeled as coming from an AI versus a human. The analysis showed that participants in the AI condition experienced significantly greater confidence drift ($M = 0.092$, $SD = 0.40$) than those in the human condition ($M = 0.074$, $SD = 0.37$), $t(34,700) = 4.46$, $p < .00001$ (Table 1). Although this result was statistically significant, the effect size was small (Cohen's $d = 0.048$). Given the large sample size, even very small effects can achieve statistical significance; therefore, this finding is unlikely to reflect a Type I error but instead indicates a statistically reliable yet practically small effect.

Table 1: Participants receiving AI-labeled advice experienced slightly higher confidence drift than those given human-labeled advice.

Advice Source	Mean Confidence Drift	Standard Deviation
AI	0.092	0.40
Human	0.074	0.37

Confidence Drift by Age Group and Advice Source:

A two-way ANOVA was conducted to examine the effects of advice source, age group, and their interaction on confidence drift. As shown in Table 2, there was a significant main effect of advice source, $F(1, 34,771) = 19.06$, $p < .00001$, and a significant main effect of age group, $F(5, 34,771) = 5.34$, $p < .0001$. The interaction between advice source and age group was marginal, $F(5, 34,771) = 2.06$, $p = .067$. These findings suggested a modest trend that different age groups may respond differently to AI versus human-labeled advice, although this effect did not reach statistical significance.

Table 2: Two-way ANOVA results show that both advice source and age group significantly influenced confidence drift, while their interaction was not statistically significant.

Source of Variation	F-value	p-value
Advice Source	19.06	< 0.00001
Age Group	5.34	< 0.0001
Advice Source x Age Group	2.06	0.067

Predicting Confidence Drift:

A multiple linear regression was conducted with age, advice source (coded 0 = human, 1 = AI), and their interaction entered as predictors of confidence drift. As shown in Table 3, age was a significant negative predictor, $\beta = -0.0011$, $p < .001$, indicating that older participants experienced smaller changes in confidence. Advice source and the interaction term

were not statistically significant. Overall, the analysis showed a very small proportion of variance ($R^2 = 0.001$), suggesting that while age was associated with confidence drift, AI labeling itself was not a strong independent predictor.

Table 3: Multiple regression results indicate that age negatively predicted confidence drift, whereas advice source did not have a significant independent effect.

Predictor	Coefficient (β)	p-value
Age	- 0.0011	< 0.001
Advice Source (AI)	+ 0.0020	0.880
Age x Advice Source	+ 0.0005	0.149

Predicting Decision Change:

A logistic regression was conducted to examine the likelihood of participants changing their answer after receiving advice. As shown in Table 4, age was a significant negative predictor, $\beta = -0.0029$, $p = .028$, indicating that older participants were less likely to change their decisions. Advice source and its interaction with age were not statistically significant, suggesting that the probability of changing answers did not differ meaningfully between AI and human-labeled advice.

Table 4: Logistic regression predicts a lower likelihood of decision change among older participants, regardless of whether the advice came from AI or human sources.

Predictor	Coefficient (β)	p-value
Intercept	+ 0.392	< 0.001
Age	- 0.0029	0.028
Advice Source (AI)	+ 0.056	0.366
Age x Advice Source	+ 0.0024	0.206

Multiple Comparisons and Statistical Limitations:

This analysis involved multiple statistical tests across two dependent variables and several interaction terms. No formal correction (e.g., Bonferroni) was applied due to the exploratory nature of the study. This approach is not without limitations, but it is consistent with methodological arguments by Rothman¹⁴ and Perneger¹⁵ that multiple-comparison adjustments can inflate Type II error rates and are inappropriate when tests address conceptually distinct hypotheses. Additionally, the primary finding in our analysis was highly significant ($p < .00001$), well below any reasonable correction threshold. Future studies should apply appropriate multiple-comparison corrections to ensure robust conclusions.

Discussion

We explored broader implications of AI-assisted decision-making through this study. Our findings support some of the earlier research and also highlight some interesting new dimensions of human-AI collaboration. Our study results confirmed that AI-labeled advice has a measurable cognitive impact on users, even though both groups received identical information. Those with AI-labeled advice expressed significantly greater confidence drift than those given human-labeled advice. Among the hypotheses tested, H1 was supported, as AI-labelled advice significantly increased confidence drift. However, H2 and H3 were not supported, as AI-labelled advice did not significantly influence decision change, and no significant interaction between age and advice source was observed. This finding is consistent with research by Zhang *et al.*,¹⁶ who found that AI framing and confidence signals

can disproportionately increase user trust without a change in performance. Our results extend those findings by demonstrating a clear dissociation between confidence changes and behavioral changes. AI-labeled advice increased participants' confidence, but there was no reciprocal change in decisions. One possible explanation for this dissociation is that AI-labeled advice may increase perceived authority or credibility, leading to greater subjective confidence without necessarily altering underlying decision processes. This may reflect cognitive anchoring, where individuals retain their initial judgments while adjusting confidence levels, or differences in how machine-generated advice is interpreted compared with human advice. Such dissociation means users feel more certain internally without necessarily making different decisions. One may characterize the irrational boost in user confidence via AI labeling as "artificial confidence". This psychological phenomenon suggests that AI advice has a limited effect on our behavior. This can have serious real-world implications with rising AI adoption. This finding is in contrast to some of the established automation bias literature. Most studies suggest that AI effects typically manifest as both cognitive and behavioral changes.^{2,4} A systematic review of automation bias by Goddard *et al.*¹⁷ also suggested that inappropriate reliance typically affects both decision-making processes and outcomes. Increased confidence without behavioral flexibility could also foster reduced openness to feedback and limit post-decision reflection in real-world applications.

We have considered several psychological mechanisms that may explain this dissociation between cognitive and decisional processes. From a cognitive perspective, AI labels may activate heuristics that enhance subjective certainty via algorithmic appreciation. This results in trusting algorithmic judgments more than human ones.¹⁸ This increased confidence is just through perceived legitimacy rather than improved understanding, which obviously doesn't result in a change in decision. AI labeling may also trigger responses through emotional pathways. Most importantly, technology optimism bias and automation complacency can create positive emotional responses to AI-generated content.¹⁹ Such an emotional boost may manifest as increased confidence without necessarily improving analytical thinking or behavioral flexibility. Both cognitive and emotional pathways may inflate confidence while leaving decision processes unchanged. This dissociation can be hypothesized through the dual-process theories of cognition.²⁰ AI labeling may primarily influence System 1 responses (fast, intuitive confidence feelings) while leaving System 2 processes (slow, deliberative decision revision) relatively unaffected. This explains why participants feel different but act similarly.

The small effect size (Cohen's $d = 0.048$) of our study needs further explanation, as even a small effect can amplify when applied to large populations or accumulate over time.²¹ In real-world scenarios, millions of users interact with AI systems daily. Small systematic confidence biases can aggregate to significant population-level effects. For instance, AI medical triage systems that modestly increase user confidence could influence help-seeking behavior across thousands of patients. Moreover, users rarely interact with AI systems only once.

Hence, small confidence increases may compound over repeated interactions, potentially leading to substantial shifts in trust and reliance patterns, as suggested by previous research.⁴ The confidence-behavior gap also raises particular concerns for AI deployment in high-stakes fields. In particular, users who feel more confident without improved decision quality may be less likely to seek additional information, consult experts, or engage in critical evaluation, resulting in potentially dangerous consequences. These findings extend beyond AI's influence on decisional accuracy alone. They demand stronger AI transparency and optimized trust calibration alongside greater AI adoption.

The analysis also revealed meaningful age-related patterns similar to prior observations by Czaja *et al.*¹¹ and Pak *et al.*¹² Age consistently predicted lower responsiveness to advice in our study. Older adults showed smaller confidence drift and lower rates of decision change. However, the interaction between age and advice source did not reach statistical significance. This suggests that the effect of AI remains relatively similar across the age groups. Younger adults may be more cognitively malleable overall, but AI-labeled advice appears to influence confidence consistently, regardless of age. Our findings urge designers to recognize AI's impact on cognition, which seems to be more pronounced than its effect on behavior. This might also be particularly important to younger users who are likely to adopt AI more readily and may therefore be more vulnerable. On the other hand, for older adults, building initial trust may be more important than simply improving algorithmic accuracy. As Rahwan *et al.*²² emphasize, understanding how AI systems influence human cognition and decision-making becomes crucial as machine behavior increasingly shapes our societies. Recent research explores these concerns in sectors like mental health and education. Zhai *et al.*²³ established that AI tools designed for mental health triage in higher education can shape student self-assessment even without any behavioral intervention. Similarly, educators appreciate AI for a personalized learning experience but express concern about students' overdependence on AI-generated feedback.²⁴ Hence, these findings may point to the potentially serious cognitive risks posed by AI-induced confidence drift in real-world learning and support environments. Age-related differences in learning may also influence this effect, as younger individuals tend to demonstrate greater cognitive flexibility, whereas older individuals may rely more on established knowledge and heuristics. This may affect their ability to critically evaluate and adapt to AI-generated advice.

An overall analysis of our results is consistent with research in the field of trust in Human-AI collaboration. For example, research by Glikson and Woolley²⁵ describes emotional and cognitive trust, which highlights subconscious emotional over-reliance on AI systems. Similar research by Alon-Barkat and Busuioc⁶ also cautions against users' selective adherence to advice based on their pre-existing biases. If we inflate confidence artificially without meaningful flexibility in behavior, users may risk losing the ability to change behavior in response to their conviction. To mitigate these risks, Srinivasan and Thomason⁵ proposed adaptive AI systems to monitor user trust and adjust

their responses through more explanation or gentle cognitive challenges. In conclusion, future AI systems should encourage confidence, critical reflection, and behavioral change across all ages and contexts.

Limitations and Future Research Directions:

As this study is based on secondary data analysis, it is limited by the structure and variables of the original dataset. The researcher had no control over the experimental design, variable selection, or measurement processes, which may restrict the scope of interpretation. Despite relying on a large dataset, several important variables are ignored. Most importantly, the dataset fails to capture measures of education, AI familiarity, digital literacy, technology anxiety, and baseline trust. These variables are likely to influence cognitive effects. The study population included a primarily English-speaking Western sample, but cultural attitudes toward technology and authority can vary significantly across societies. Hence, results are not generalizable across cultures. The analysis also examined single-session interactions. Therefore, it is not representative of real-world AI use, which involves repeated exposure over time. Future longitudinal study designs should examine the evolution of confidence effects and determine if minor confidence enhancements compound into significant trust miscalibration over a period. Future studies should incorporate multiple comparison corrections and experimental designs that can manipulate AI labeling while controlling for actual advice quality. This can strengthen causal inferences about labeling effects and improve generalizability.

Conclusion

Our study shows that AI-labeled advice increases user confidence modestly compared to human-labeled advice, despite both providing the same information. But the cognitive process of a rise in confidence does not lead to greater behavioral flexibility or changes in decision-making. The relative effect of AI versus human advice remained consistent across all age groups, even though age came out as an important negative predictor of shift in confidence and change in decisions. These findings raise concerns about the growing use of AI-based advice in areas where decision accuracy is critical, such as healthcare, education, and public policy. Finally, the findings of this study highlight the need for AI systems designed to build trust while still encouraging critical reflection and independent judgment.

Acknowledgments

I would like to thank Dr. Naveen Thomas, Director of Clinical Services at Sunshine Hospital, Western Health, for his valuable support and guidance throughout this research. I am particularly grateful for his assistance in understanding the statistical methods and for providing insightful ideas that contributed to this study.

References

- Hancock, P. A.; Billings, D. R.; Schaefer, K. E.; Chen, J. Y. C.; de Visser, E. J.; Parasuraman, R. How and why humans trust: A meta-analysis and elaborated model. *Front. Psychol.* **2023**, *14*, 1094626.
- Dzindolet, M. T.; Peterson, S. A.; Pomranky, R. A.; Pierce, L. G.; Beck, H. P. The role of trust in automation reliance. *Int. J. Hum.-Comput. Stud.* **2003**, *58* (6), 697–718.
- Jacovi, A.; Marasović, A.; Miller, T.; Goldberg, Y. Formalizing trust in artificial intelligence: Prerequisites, causes and goals of human trust in AI. In *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency (FAccT '21)*; Association for Computing Machinery: New York, NY, **2021**; pp 624–635.
- Parasuraman, R.; Riley, V. Humans and automation: Use, misuse, disuse, abuse. *Hum. Factors* **1997**, *39* (2), 230–253.
- Srinivasan, T.; Thomason, J. Adjust for trust: Mitigating trust-induced inappropriate reliance on AI assistance. *arXiv preprint arXiv:2502.13321* **2025**.
- Alon-Barkat, S.; Busuioc, M. Human-AI interactions in public sector decision making: "Automation bias" and "selective adherence" to algorithmic advice. *J. Public Adm. Res. Theory* **2023**, *33* (1), 153–169.
- Yaniv, I.; Kleinberger, E. Advice taking in decision making: Ego-centric discounting and reputation formation. *Organ. Behav. Hum. Decis. Process.* **2000**, *83* (2), 260–281.
- Bonaccio, S.; Dalal, R. S. Advice taking and decision-making: An integrative literature review, and implications for the organizational sciences. *Organ. Behav. Hum. Decis. Process.* **2006**, *101* (2), 127–151.
- Vodrahalli, K.; Daneshjou, R.; Gerstenberg, T.; Zou, J. Do humans trust advice more if it comes from AI? An analysis of human-AI interactions. In *Proceedings of the 2022 AAAI/ACM Conference on AI, Ethics, and Society*; ACM: New York, NY, **2022**; pp 763–777.
- Hoff, K. A.; Bashir, M. Trust in automation: Integrating empirical evidence on factors that influence trust. *Hum. Factors* **2015**, *57* (3), 407–434.
- Czaja, S. J.; Charness, N.; Fisk, A. D.; Hertzog, C.; Nair, S. N.; Rogers, W. A.; Sharit, J. Factors predicting the use of technology: Findings from the Center for Research and Education on Aging and Technology Enhancement (CREATE). *Psychol. Aging* **2006**, *21* (2), 333–352.
- Pak, R.; Rovira, E.; McLaughlin, A. C.; Baldwin, N. Does the domain of technology impact user trust? Investigating trust in automated systems for public safety, military, and consumer domains. *Hum. Factors* **2017**, *59* (4), 582–592.
- Buçinca, Z.; Malaya, M. B.; Gajos, K. Z. To trust or to think: Cognitive forcing functions can reduce overreliance on AI in AI-assisted decision-making. *Proc. ACM Hum.-Comput. Interact.* **2021**, *5* (CSCW1), 1–21.
- Rothman, K. J. No adjustments are needed for multiple comparisons. *Epidemiology* **1990**, *1* (1), 43–46.
- Perneger, T. V. What's wrong with Bonferroni adjustments. *BMJ* **1998**, *316* (7139), 1236–1238.
- Zhang, Y.; Liao, Q. V.; Bellamy, R. K. E. Effect of confidence and explanation on accuracy and trust calibration in AI-assisted decision making. In *Proceedings of the 2020 Conference on Fairness, Accountability, and Transparency*; ACM: New York, NY, **2020**; pp 295–305.
- Goddard, K.; Roudsari, A.; Wyatt, J. C. Automation bias: A systematic review of frequency, effect mediators, and mitigators. *J. Am. Med. Inform. Assoc.* **2012**, *19* (1), 121–127.
- Logg, J. M.; Minson, J. A.; Moore, D. A. Algorithm appreciation: People prefer algorithmic to human judgment. *Organ. Behav. Hum. Decis. Process.* **2019**, *151*, 90–103.
- Lee, J. D.; See, K. A. Trust in automation: Designing for appropriate reliance. *Hum. Factors* **2004**, *46* (1), 50–80.
- Kahneman, D. *Thinking, fast and slow*; Farrar, Straus and Giroux: New York, **2011**.

21. Funder, D. C.; Ozer, D. J. Evaluating effect size in psychological research: Sense and nonsense. *Adv. Methods Pract. Psychol. Sci.* **2019**, *2* (2), 156–168.
22. Rahwan, I.; Cebrian, M.; Obradovich, N.; Bongard, J.; Bonnefon, J. F.; Breazeal, C.; Crandall, J. W.; Christakis, N. A.; Couzin, I. D.; Jackson, M. O.; Jennings, N. R.; Kamar, E.; Kloumann, I. M.; Larson, K.; Lazer, D.; McElreath, R.; Mislove, A.; Parkes, D. C.; Pentland, A. S.; Roberts, M. E.; Shariff, A.; Tenenbaum, J. B.; Wellman, M. Machine behavior. *Nature* **2019**, *568* (7753), 477–486.
23. Zhai, S.; Zhang, S.; Rong, Y.; Rong, G. Technology-driven support: Exploring the impact of artificial intelligence on mental health in higher education. *Educ. Inf. Technol.* **2025**, *30*, 13931–13959.
24. Delello, J. A.; Sung, W.; Mokhtari, K.; Hebert, J.; Bronson, A.; De Giuseppe, T. AI in the classroom: Insights from educators on usage, challenges, and mental health. *Educ. Sci.* **2025**, *15* (2), 113.
25. Glikson, E.; Woolley, A. W. Human trust in artificial intelligence: Review of empirical research. *Acad. Manage. Ann.* **2020**, *14* (2), 627–660.

■ Authors

Kanishk Mahasuar is a Year 10 high school student from Australia with a strong passion for mathematics and science. He is interested in exploring innovative technologies and their applications to real-world problems. Beyond his academic pursuits, Kanishk enjoys rowing, debating, and playing competitive chess.