

Human Agency in AI Use

Exploratory findings from surveyed generative AI users

NEURAVOX POLICY LAB

RESEARCH REPORT, AUGUST 2026

Gideon Abako

Neuravox Foundation

neuravox.org

About this report

This report presents the first exploratory findings from the Neuravox Human Agency Survey. The study examines how people use generative AI and how that use relates to independent thinking, verification, explanation, capability retention and decision authority. The report combines primary survey evidence with a review of research and policy sources on cognitive offloading, automation bias, critical thinking, AI literacy and human autonomy.

The study is part of the Neuravox Policy Lab research programme on human agency in everyday AI use. Its purpose is to build an evidence base for a practical Human Agency Framework that can be tested and refined in later research.

Suggested citation: Abako, G. L. (2026). Human Agency in AI Use: Exploratory findings from surveyed generative AI users. Neuravox Policy Lab, Neuravox Foundation.

Acknowledgements: Neuravox Foundation thanks everyone who participated in the survey and everyone who shared it through professional and online communities.

Contents

Executive summary

1. Why human agency matters in everyday AI use
2. What existing evidence tells us
3. Study design and analysis
4. Who took part and how AI was used
5. What participants reported about human agency
6. Perceived change in independent thinking
7. Relationships across agency, AI use and perceived change
8. What the findings mean
9. Emerging Human Agency Framework
10. Implications for practice and policy
11. Research agenda
12. Scope of interpretation
13. Conclusion

References

Appendices

Executive summary

Generative AI can now draft, summarize, reason, plan, code and recommend. These capabilities make AI useful across everyday work and life. They also move part of the thinking process outside the user. The central question for this study is how people can gain value from AI while retaining control over judgment, understanding, skill and important decisions.

Neuravox Foundation surveyed 76 adults who had used generative AI for at least three months. Most participants were experienced users. Nearly seven in ten used AI daily or several times a day, and more than three quarters had used it for at least one year. The survey measured four forms of delegation and three forms of retained agency, then asked participants how their ability to think through difficult tasks without AI had changed since they began using generative AI regularly.

Frequent use did not signal weaker agency.

AI use frequency had no meaningful relationship with the Delegation score or perceived change in independent thinking. Frequency showed a modest positive relationship with the Retention score. In this sample, how people used AI carried more information than how often they used it.

Retained control was common alongside moderate delegation.

The mean Delegation score was 2.71 out of 5 and the mean Retention score was 4.02. Eighty percent of respondents said they could explain their final work often or almost always. Seventy eight percent said they often or almost always questioned AI when it conflicted with their judgment.

Perceived outcomes were mixed.

Forty two respondents, 55.3 percent, reported improvement in their ability to think through difficult tasks without AI. Twenty respondents, 26.3 percent, reported some reduction. Nine reported no noticeable change and five were unsure. The pattern supports continued study of both benefit and dependence.

Human agency appears to have several distinct dimensions.

The seven survey items did not support one overall agency score. The findings point toward separate capabilities involving independent initiation, cognitive authorship, verification, explanatory ownership, capability retention and decision authority.

The strongest implication is straightforward. Usage intensity is a poor shortcut for human agency. A person may use AI many times each day while still checking outputs, challenging recommendations, understanding the final work and retaining the ability to function without the system. A useful Human Agency Framework therefore needs to examine the distribution of cognitive and decisional authority between the person and AI.

The report develops six working domains from the survey findings and existing evidence. These domains will guide the next phase of Neuravox research, including qualitative interviews, improved measurement and a larger validation study.

1. Why human agency matters in everyday AI use

Generative AI has moved quickly from a specialist technology into ordinary work, study and personal activity. People use systems such as ChatGPT, Claude, Gemini and Copilot to search for information, draft text, analyse data, write code, plan projects, learn new material and think through personal questions. Many of these tasks involve more than execution. They involve deciding what matters, how a problem should be framed, which evidence should be trusted and what conclusion should follow.

Human agency concerns the ability to direct one's own thinking and action. In the context of generative AI, this includes the capacity to begin a task with an independent view, shape the structure of ideas, verify information, challenge an AI response, explain the final result, retain relevant capability and keep authority over important decisions. Agency can therefore change even when the visible output remains good.

This creates a measurement problem. Counting prompts, hours of use or frequency of access tells us how much AI is present in a person's workflow. It says little about who is steering the reasoning. A frequent user may treat AI as a source of options and still exercise strong judgment. A less frequent user may accept a complete answer with minimal review. The study was designed around this distinction.

The research question guiding the wider Neuravox programme is: How can people retain cognitive, decisional and social agency while using generative AI as part of everyday work and life? This first survey focuses on cognitive and decisional behaviours that can be reported by users in a short exploratory instrument.

2. What existing evidence tells us

2.1 Cognitive offloading is a normal part of human thinking

People have always used external tools to reduce mental effort. Notes, calculators, calendars, search engines and shared records all move part of a cognitive task into the environment. Risko and Gilbert (2016) describe this process as cognitive offloading and show that the decision to offload depends partly on task demands and on people's own judgments about their mental ability. Offloading can be useful. Its effects depend on the task, the tool and the way the person remains engaged with the work.

Research on automation adds a second concern. People can give automated recommendations more weight than they deserve, especially when verification is difficult or attention is constrained. Reviews of automation bias show that trust, task complexity, workload and the design of decision support systems can affect the degree of reliance on automated advice (Parasuraman & Manzey, 2010; Goddard, Roudsari, & Wyatt, 2012; Romeo & Conti, 2026). These findings make verification and decision authority important parts of any framework for human agency.

2.2 Generative AI expands the range of cognition that can be delegated

Generative AI can take on reasoning, synthesis, drafting, evaluation and creative generation. This expands cognitive offloading beyond memory or calculation. Lee and colleagues (2025) surveyed 319 knowledge workers and collected 936 examples of generative AI use in work tasks. They found that confidence in AI was associated with less critical thinking, while confidence in one's own ability was associated with more critical thinking. Their qualitative findings also

showed that critical thinking changes shape in AI assisted work. Verification, response integration and task stewardship become more important.

A 2026 study by Zhu and colleagues provides an especially close comparison to the Neuravox study. Their three wave survey of 589 university students and early career knowledge workers separated dependent cognitive offloading from autonomous cognitive offloading. Dependent use involved allowing AI to substitute for core thinking. Autonomous use involved using AI as support while the person retained control of the cognitive process. The two patterns had different associations with perceived downstream capability, motivation, deep processing and independent judgment. The authors conclude that the manner of AI engagement deserves attention alongside frequency of use (Zhu et al., 2026).

2.3 Agency, autonomy and literacy are multidimensional

Recent work on AI autonomy reaches a similar conclusion from a different direction. Buijsman, Carter and Bermudez (2025) distinguish skilled competence from authentic value formation in AI decision support and propose design practices that preserve reflective judgment. HumanAgencyBench takes a system focused approach and evaluates whether AI assistants support user agency through behaviours such as correcting misinformation, deferring important decisions, encouraging learning and maintaining social boundaries (Sturgeon et al., 2025). These approaches show that agency includes several capabilities and relationships between a person and a system.

Measurement work on AI literacy also offers a useful lesson. MAILS, GLAT and SAIL4ALL use multiple dimensions or separate components to assess knowledge, skills, self management and critical evaluation (Carolus et al., 2023; Jin et al., 2024; Soto-Sanfiel, Angulo-Brunet, & Lutz, 2025). GLAT also shows the value of testing actual knowledge alongside self report. SAIL4ALL explicitly advises against combining all literacy themes into one overall score because the construct is multidimensional. Human agency is distinct from AI literacy, but the measurement lesson is relevant. Complex human capabilities require evidence before they are compressed into a single index.

2.4 Policy frameworks already place human determination at the centre

International AI governance frameworks provide a normative foundation for this work. UNESCO's Recommendation on the Ethics of Artificial Intelligence states that AI should not displace ultimate human responsibility and calls for research on the effects of AI based recommendations on decision autonomy (UNESCO, 2021). The OECD AI Principles updated in 2024, place autonomy, human agency and oversight within the principle on human rights and democratic values (OECD, 2024).

The African Union Continental Artificial Intelligence Strategy also adopts a people centred approach and links responsible AI to human rights, dignity, inclusion and African development priorities (African Union, 2024). These frameworks establish the importance of human oversight. The practical challenge is to translate that principle into behaviours that can be observed in everyday AI use. The Neuravox study contributes to that operational question.

3. Study design and analysis

3.1 Study design

The Neuravox Human Agency Survey used an exploratory cross sectional design. Data were collected online from 8 August to 22 August 2026. Recruitment used professional networks,

LinkedIn and online communities including Slack groups. Participants were encouraged to share the survey with other generative AI users. This produced a convenience and snowball sample suited to exploratory analysis.

Participation was open to adults aged 18 or older who had used generative AI for at least three months. Respondents saw a short study information statement and provided consent before the survey questions became available. The survey did not request names, email addresses, phone numbers or employers.

3.2 Survey measures

The survey recorded the respondent's main current activity, frequency of generative AI use, length of experience and main uses of AI. Main uses could include work, study, writing, research, coding, planning, personal advice and other activities.

Seven behavioural items used a five point response scale from Never to Almost always. Four items captured delegation: asking AI before trying to think independently, using AI output without checking, allowing AI to structure ideas, and delegating decisions previously made personally. Three items captured retained agency by questioning AI when it conflicted with one's judgment, being able to explain the final work, and feeling confident completing similar tasks without AI.

The final survey question asked how the respondent's ability to think through difficult tasks without AI had changed since they began using generative AI regularly. Response options ranged from Reduced a lot to Improved a lot, with a separate Not sure category.

3.3 Analysis

The analysis used all 76 complete and consenting responses. Likert items were coded from 1 to 5. A Delegation score was calculated as the mean of the four delegation items. A Retention score was calculated as the mean of the three retained agency items. A candidate overall agency score was tested after reversing the delegation items. Its internal consistency was too weak for use in the primary analysis, so the report keeps the individual behaviours visible and treats Delegation and Retention as exploratory summaries.

The analysis used descriptive statistics, Spearman correlations, bootstrap confidence intervals, Kruskal Wallis tests for grouped comparisons, and Mann Whitney tests for exploratory comparisons across common AI use cases. Benjamini Hochberg correction was used for families of multiple tests. A restrained ordinal regression was used as a sensitivity analysis. The full analysis was reproducible from the frozen dataset and passed the project validation checks.

3.4 Interpretation

The study describes the 76 people who took part. It does not estimate the prevalence of these behaviours in a wider population. The survey captures self reported behaviour at one point in time, so associations are interpreted as patterns for further study. The findings are used to develop research questions and working dimensions for the Human Agency Framework.

4. Who took part and how AI was used

The sample included people in employment, freelancing, self employment, academic work, study and periods outside paid work. The largest group was employed respondents, with 25 participants. Freelancers or consultants accounted for 13 participants. Students and people not currently working each accounted for 10. Self employed business owners and academics or researchers each accounted for nine.

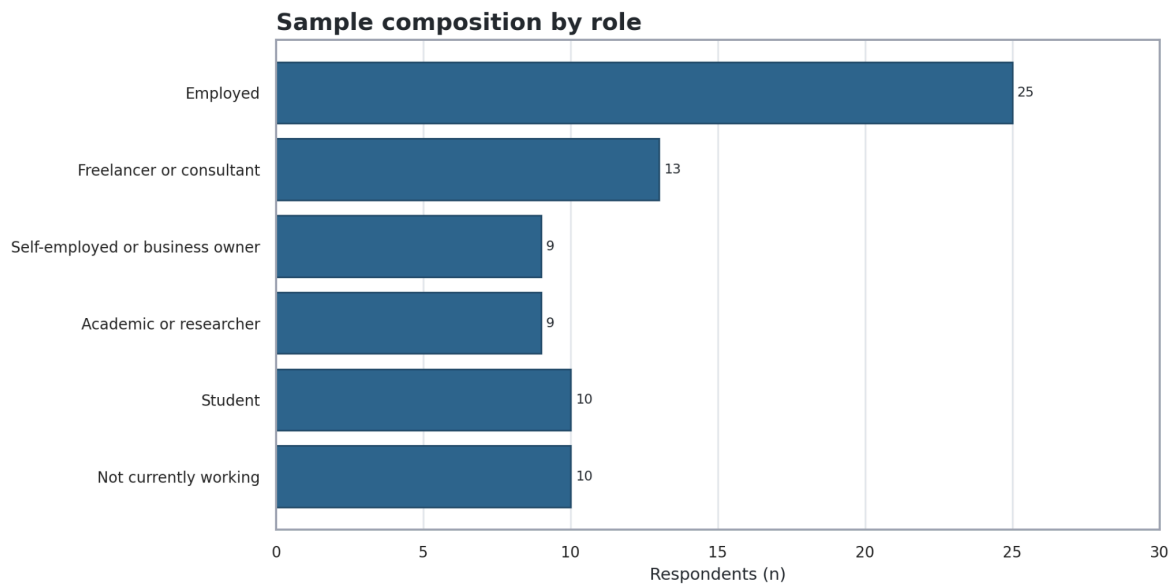


Figure 1. Sample composition by main current activity. Counts are shown for all 76 respondents.

AI use was frequent and established for most respondents. Thirty participants used generative AI daily and 23 used it several times a day. Together, 53 participants, 69.7 percent of the sample used AI at least daily. Thirty two respondents had more than two years of experience and 27 had one to two years. This means 59 respondents, 77.6 percent had used generative AI for at least one year.

AI use intensity and experience

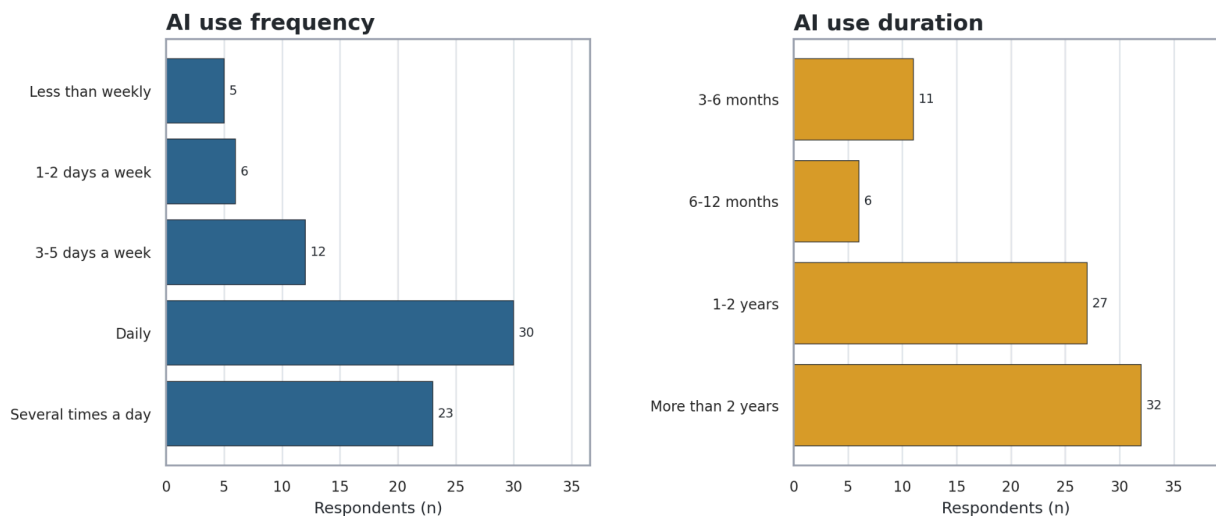


Figure 2. Frequency and duration of generative AI use among the 76 respondents.

Writing or editing was the most common use, selected by 53 respondents. Research or information was selected by 52, study or learning by 49 and work tasks by 46. Planning or brainstorming was selected by half the sample. Twenty nine respondents used AI for personal advice or decisions and 22 used it for coding or data analysis. Respondents could select several uses so these categories overlap.

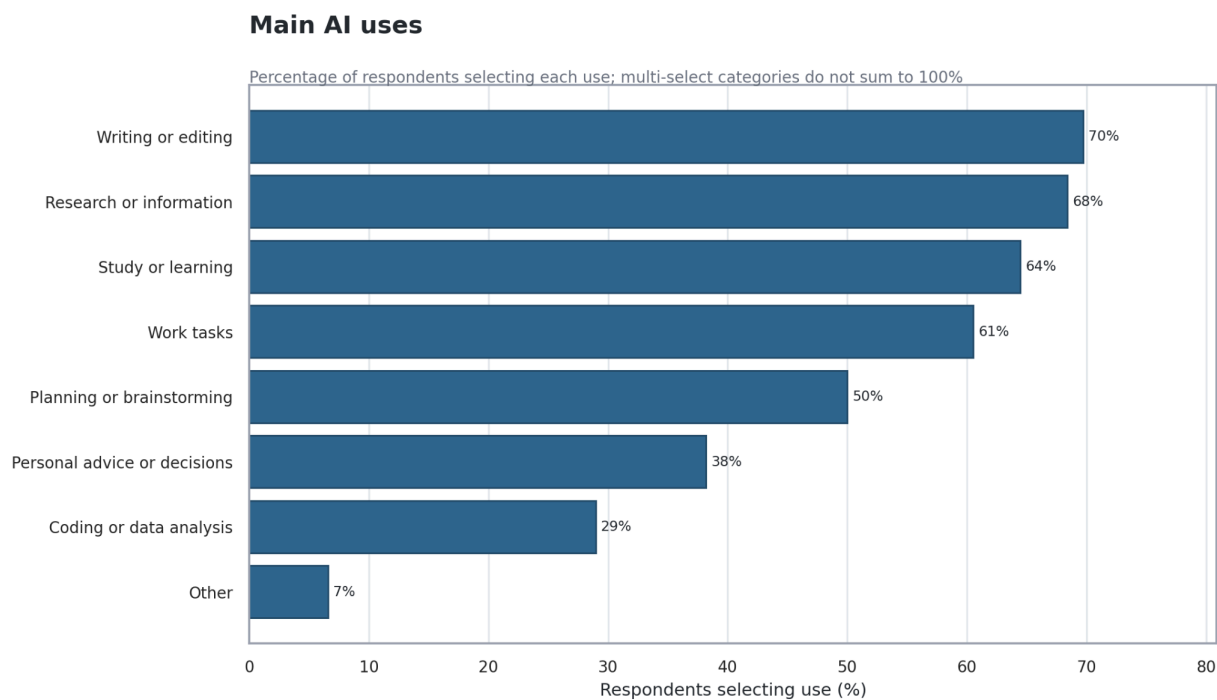


Figure 3. Main uses of generative AI. Respondents could select more than one use.

5. What participants reported about human agency

The seven behavioural questions produced a clear descriptive pattern. Delegation was present but most delegation items centred on Sometimes. Retained agency behaviours were more strongly concentrated in Often and Almost always. The combination matters because it shows that use of AI assistance and retention of human control can appear together in the same sample.

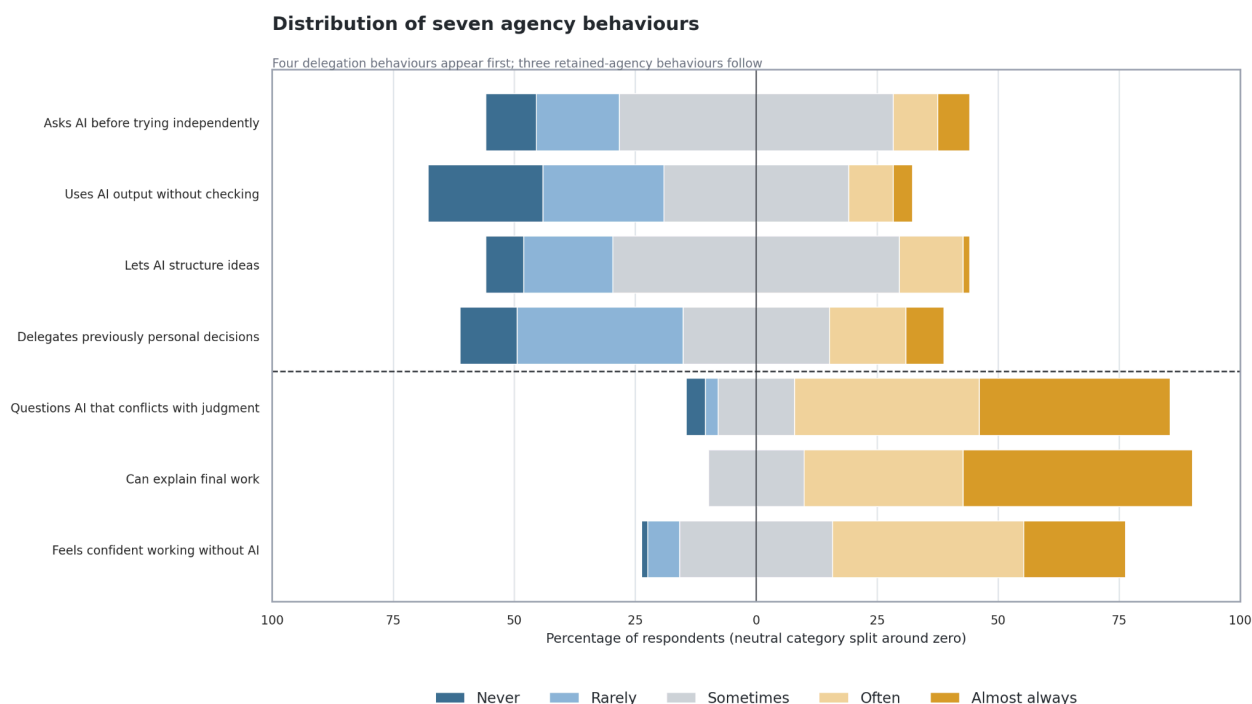


Figure 4. Distribution of the seven agency related behaviours. The four delegation items appear first, followed by the three retained agency items.

5.1 Delegation behaviours

Asking AI before trying to think independently had a mean score of 2.84 out of 5. Forty three respondents, 56.6 percent selected Sometimes. Twelve selected Often or Almost always. This suggests that AI often enters the reasoning process early but habitual first response delegation was concentrated in a smaller group.

Using AI output without checking had the lowest delegation mean at 2.45. Thirty seven respondents, 48.7 percent selected Never or Rarely. Twenty nine selected Sometimes. Ten selected Often or Almost always. Verification therefore appears to be a common part of AI use for many respondents while a meaningful minority reported regular use of unchecked output.

Allowing AI to structure ideas had a mean of 2.82. Forty five respondents, 59.2 percent selected Sometimes. Eleven selected Often or Almost always. This item captures cognitive authorship because structuring a task or argument determines how later reasoning is organised.

Delegating decisions previously made personally had a mean of 2.74. Thirty five respondents, 46.1 percent selected Never or Rarely while 18, 23.7 percent selected Often or Almost always. This behaviour becomes especially important when AI use moves from drafting and information work into choices with personal consequences.

5.2 Retained agency behaviours

Questioning AI when it conflicted with one's judgment was common. Fifty nine respondents, 77.6 percent selected Often or Almost always. The mean was 4.07. This indicates a strong reported willingness to challenge AI output when users detect a conflict with their own knowledge or belief.

The strongest retained agency behaviour was the ability to explain final work. Sixty one respondents, 80.3 percent selected Often or Almost always and no respondent selected Never or Rarely. The mean was 4.28. This suggests that many participants saw themselves as retaining understanding and explanatory ownership even when AI contributed to the task.

Confidence in completing similar tasks without AI was also relatively high, with a mean of 3.72. Forty six respondents, 60.5 percent selected Often or Almost always. Thirty selected Sometimes, Rarely or Never. Capability retention therefore appears more variable than explanation and conflict questioning.

6. Perceived change in independent thinking

The survey asked respondents to assess how their ability to think through difficult tasks without AI had changed since they began using generative AI regularly. Forty two respondents, 55.3 percent reported improvement. Twenty seven said Improved a lot and 15 said Improved somewhat.

Twenty respondents, 26.3 percent reported some reduction. Fifteen selected Reduced somewhat and five selected Reduced a lot. Nine reported no noticeable change and five selected Not sure. The distribution gives the study an important balance. Participants described both gains and losses which supports a framework that can recognise supportive and dependent patterns of AI use.

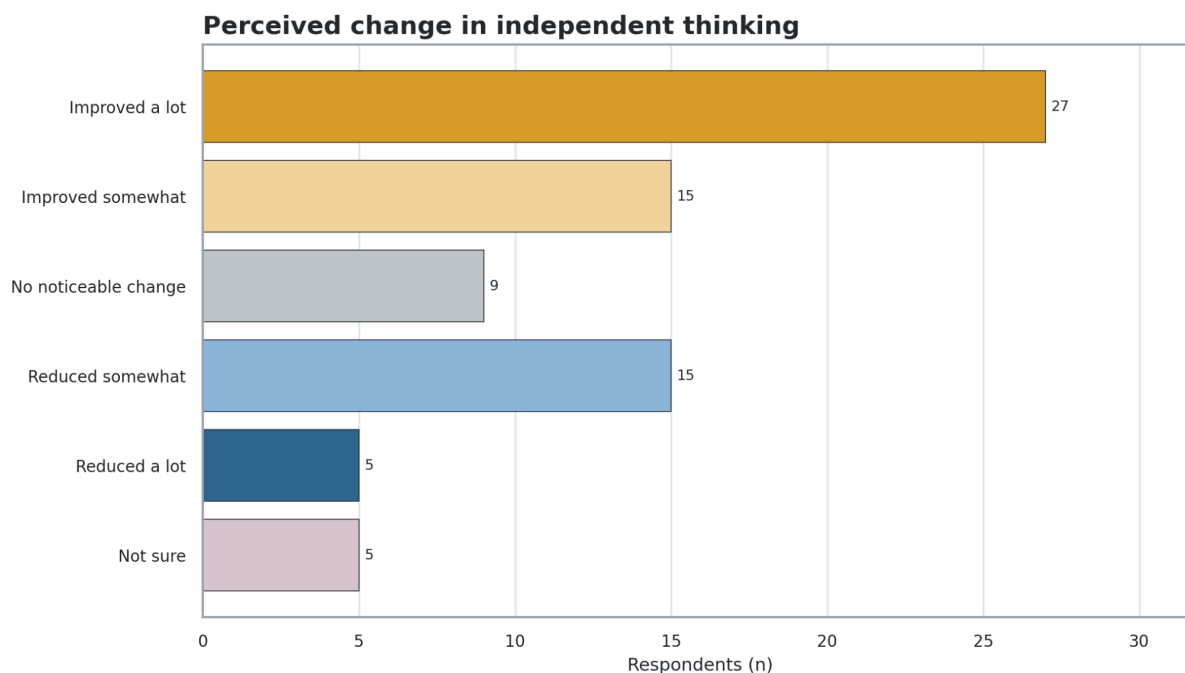


Figure 5. Perceived change in the ability to think through difficult tasks without AI.

The meaning of improvement needs careful interpretation. A participant may understand improvement as becoming more effective at solving difficult problems, becoming better at using AI or becoming more capable without AI. The survey question was designed around independent thinking but it did not ask respondents to explain the meaning behind their answer. This becomes a priority for the qualitative phase.

7. Relationships across agency, AI use and perceived change

7.1 The survey did not support one overall agency score

The mean Delegation score was 2.71 and the mean Retention score was 4.02. The candidate seven item overall score had Cronbach's alpha of 0.575 and failed the analysis rule for a coherent composite. Delegation had alpha of 0.641. Retention had alpha of 0.410. These results support continued use of the individual behaviours and caution against presenting human agency as one validated scale at this stage.

The weak Retention consistency is substantively informative. Questioning AI, explaining final work and feeling capable without AI may all contribute to agency but they represent different human capacities. A person can be strong on one and weaker on another. The future framework should preserve that distinction.

7.2 Frequent use was not associated with greater delegation

AI use frequency had almost no association with Delegation, with Spearman ρ of 0.046 and p of 0.692. Frequency was also unrelated to perceived change in independent thinking, with ρ of 0.022 and p of 0.855. Frequency had a modest positive association with Retention with ρ of 0.247, a 95 percent bootstrap confidence interval from 0.017 to 0.466, and p of 0.031.

Length of experience with AI showed no association with Delegation, Retention or perceived change. The data therefore separate exposure from agency. High use intensity can coexist with frequent checking, conflict questioning, explanation and confidence without AI.

7.3 Delegation was unrelated to perceived improvement

Delegation had essentially no relationship with perceived change, with rho of 0.010 and p of 0.934. Retention showed a modest negative association with perceived improvement, with rho of -0.257 and p of 0.031. The grouped comparison across Reduced, No change and Improved categories did not show a difference in either Delegation or Retention scores.

Agency dimensions by perceived-change group

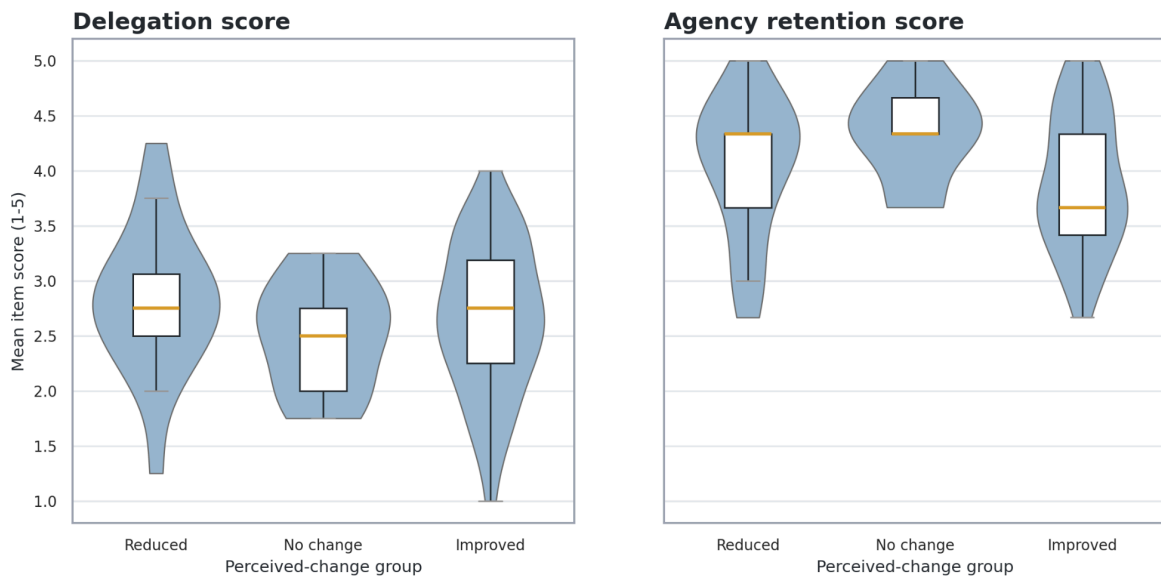


Figure 6. Delegation and Retention scores across respondents who reported reduced, unchanged or improved independent thinking. Not sure responses are excluded from this comparison.

The negative Retention relationship deserves investigation not a quick explanation. One possibility is that people with strong independent capability had less room to perceive improvement. Another is that respondents interpreted improvement as better performance with AI support. Sampling variation is also possible. The current data do not distinguish among these explanations.

7.4 Delegation and retained agency can coexist

The scatterplot of Delegation and Retention shows respondents across a broad range of combinations. Some frequent AI users report substantial delegation and still report strong retained agency. Others show low delegation with more moderate retention. The relationship between the two scores is modest, which reinforces the idea that they capture different aspects of AI use.

Delegation and retained agency can coexist

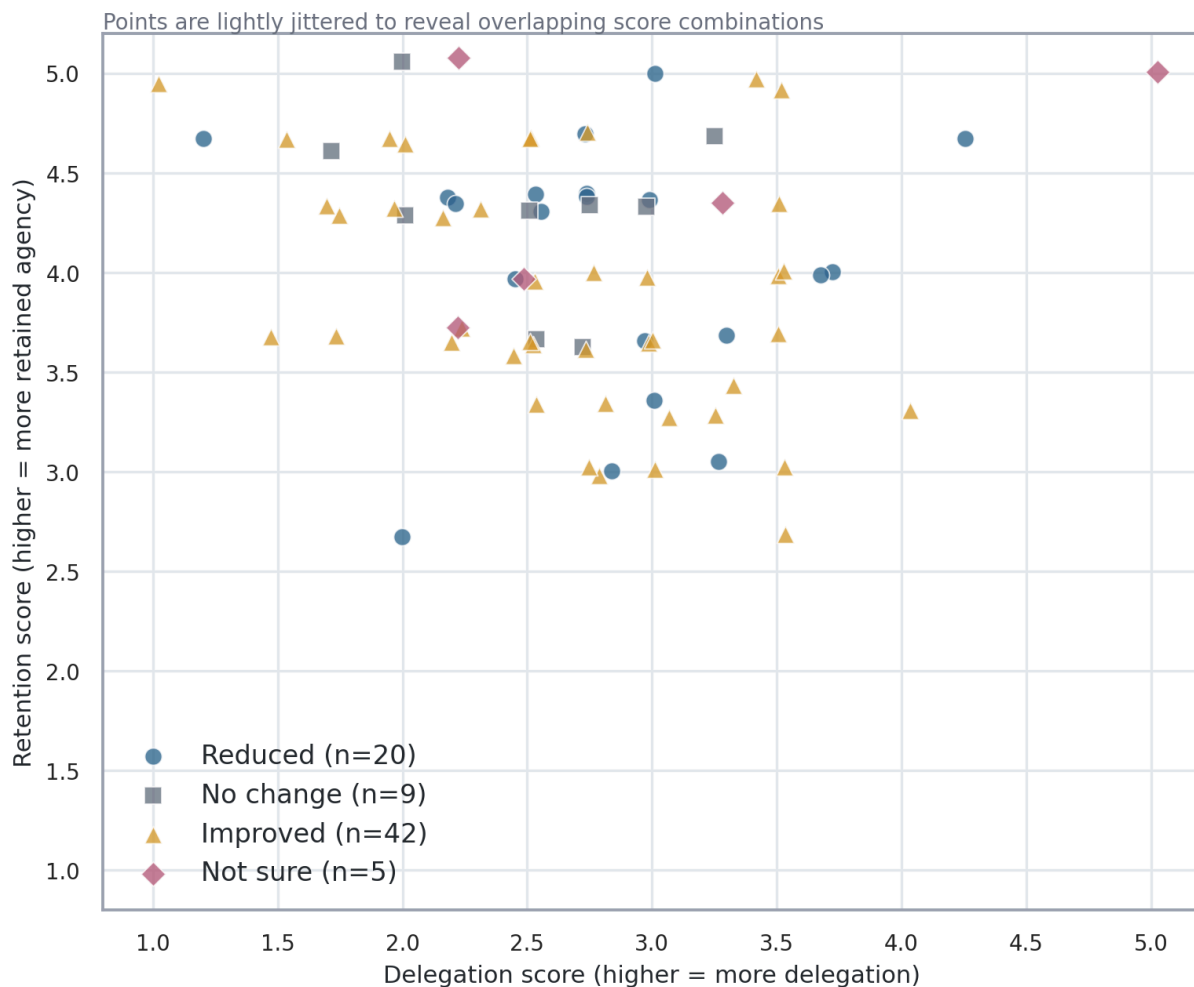


Figure 7. Delegation and Retention scores by perceived change category. The figure is descriptive and does not define respondent types or clusters.

7.5 Use case patterns point to questions for later research

Exploratory comparisons suggested higher Retention among respondents who used AI for coding or data analysis and among those who used it for research or information. The adjusted p values were 0.051 and 0.068, so these patterns did not cross the study's multiple testing threshold. Personal advice or decision users showed a moderate raw difference toward higher Delegation, but that result also weakened after correction.

These patterns are useful as hypotheses. Coding and research often provide external ways to test an AI response through executable results, source checks or comparison with evidence. Personal advice may offer fewer objective verification signals. The next phase can examine whether the availability of external checks affects how people retain agency.

8. What the findings mean

8.1 Usage intensity is context, not an agency diagnosis

The clearest empirical result is the separation between frequency and agency. Daily or several times daily use was common but frequency did not track Delegation or perceived decline. This aligns with Zhu and colleagues' distinction between dependent and autonomous offloading. The

quality of engagement matters because the same amount of AI use can involve very different distributions of cognitive control.

This finding has practical consequences. Policies that use hours, prompts or frequency as a direct proxy for human dependence will miss critical differences in verification, explanation, retained capability and decision authority. Usage intensity remains useful as contextual information. It should be interpreted alongside behaviour.

8.2 Delegation can be productive when authority remains visible

Delegation is part of why AI is useful. People use technology to reduce effort, accelerate routine work and access capabilities they do not need to reproduce manually every time. The agency question concerns the location of control. A person can delegate drafting or idea generation and still decide what matters, verify the output, explain the result and make the final choice.

The present data support this distinction because moderate delegation appears alongside strong reported explanation and conflict questioning. This creates space for a framework that evaluates the quality of delegation instead of treating every transfer of work to AI as a loss of agency.

8.3 Verification and explanatory ownership deserve central roles

Verification and explanation stand out in both the literature and the survey. Automation bias research shows why unchecked recommendations can create errors. Lee and colleagues describe verification and response integration as central parts of critical thinking in AI assisted knowledge work. In the survey, questioning conflicting AI and explaining final work were among the strongest retained agency behaviours.

These behaviours are useful because they can be translated into practice. A team can ask whether people check important claims. A reviewer can ask whether the user understands the final output. An AI product can make sources, uncertainty and decision points easier to inspect. These are observable behaviours that move human oversight from principle to practice.

8.4 Capability retention is a separate question

Understanding an AI assisted result does not guarantee that the person can complete a similar task independently. The survey reflects this distinction. Ability to explain final work was stronger than confidence without AI. Capability retention therefore deserves its own domain in the framework and should eventually be measured with task based evidence, not confidence alone.

8.5 Perceived benefit and preserved agency can move differently

More than half of respondents perceived improvement, yet the Retention score did not rise with perceived improvement. This suggests that benefit and agency should be measured separately. AI can improve speed, output quality or access to expertise while changing how much thinking the person performs. A strong framework should capture both sides without assuming that one automatically determines the other.

9. Emerging Human Agency Framework

The survey and evidence review support a multidimensional starting point for the Neuravox Human Agency Framework. The domains below are working domains. They are grounded in the current study and literature and they will be refined through qualitative interviews and later validation research.

Working domain	Meaning	Core question
Independent initiation	The person forms an initial view, question or attempt before handing the core problem to AI.	Who starts the thinking?
Cognitive authorship	The person retains influence over how ideas, arguments and tasks are structured.	Who shapes the reasoning?
Verification and epistemic judgment	The person checks important outputs, notices conflicts and can reject an AI answer.	Who decides what is credible?
Explanatory ownership	The person can explain the reasoning and final result in their own terms.	Who understands and stands behind the output?
Capability retention	The person retains enough knowledge and skill to perform relevant tasks when AI is unavailable or unsuitable.	What can the person still do independently?
Decision authority	The person retains final authority over choices that carry personal, professional or social consequences.	Who makes the consequential choice?

Two additional layers sit around these domains. The first is context of use including frequency, experience, task type, stakes, reversibility and the availability of external verification. The second is outcome including perceived effectiveness, learning, confidence, dependence and change in independent capability. Keeping these layers separate reduces the risk of treating an outcome such as speed or satisfaction as direct evidence of preserved agency.

The current study also suggests that a future framework should report a profile across domains. A single total score may become useful after measurement work but the first version should show where agency is strong and where it may be shifting.

10. Implications for practice and policy

10.1 For individuals

People can use a simple set of questions to protect agency during everyday AI use. Before an important task, ask whether you have formed your own view of the problem. During the task, check claims that matter and notice where AI is shaping the structure of your reasoning. Before using the result, make sure you can explain it and identify which decisions remain yours. For recurring tasks, periodically test whether you can still perform the essential parts without AI support.

10.2 For organizations

Organizations can move beyond generic rules about how often staff use AI. Governance can focus on the tasks where independent judgment matters, the points where verification is required and the decisions that need named human responsibility. Review processes can ask whether staff can explain AI assisted work and whether important claims were checked against reliable sources or domain knowledge.

Capability retention also deserves attention in roles where people may need to act during system failure, emergencies or unusual cases. Periodic practice without AI can reveal where essential knowledge has weakened and where AI is functioning as useful support.

10.3 For AI designers

AI products can support agency by making it easier for users to inspect sources, compare options, see uncertainty, challenge recommendations and understand how an answer was formed. Systems can encourage users to provide their own goals or initial reasoning before generating a complete solution in tasks where learning or judgment matters. Decision support can also make the boundary between advice and final authority explicit.

10.4 For policy and public interest technology

UNESCO, OECD and African Union frameworks already establish human oversight, dignity and responsibility as core principles. The next step is operational. Agency can be assessed through concrete behaviours such as verification, explanation, capability retention and decision authority. This creates a bridge between high level governance principles and the everyday experience of people using generative AI.

11. Research agenda

This first study is designed to generate the next questions. The quantitative results now give the qualitative phase a clear focus. Neuravox plans to use interviews to understand why people delegate, what they mean when they report improvement, how they decide when to trust AI, and when AI begins to change their confidence or independent capability.

A focused interview study with approximately 10 to 15 generative AI users can explore several issues that the survey could not explain. These include the difference between using AI as a starting point and using it as the primary source of reasoning, the experience of disagreeing with AI, how people verify outputs in different task domains and what happens when AI becomes unavailable.

The measurement work should then add several indicators for each proposed domain. Future studies can combine self report with task based measures. Verification can be observed through source checking. Explanatory ownership can be assessed by asking people to defend an AI assisted answer without reopening the chat. Capability retention can be tested through comparable tasks completed with and without AI. Decision authority can be studied through realistic scenarios with different levels of consequence.

A larger validation study can then test the structure, reliability and generalizability of the framework across populations and use cases. Longitudinal work will be especially valuable for understanding whether repeated patterns of delegation change capability over time.

12. Scope of interpretation

The study uses 76 respondents recruited through professional networks and online communities. The findings describe this group and provide an exploratory evidence base for framework development. Geography and demographic characteristics were not collected, so the report makes no claims about national, regional or demographic populations.

The measures are self reported and were collected at one point in time. They describe how respondents understand their own behaviour and perceived change. Objective performance and long term capability require separate measurement. The association analyses therefore identify patterns that can guide later research.

The Delegation and Retention scores are compact summaries of the survey items. Their reliability results show that the underlying behaviours should remain visible. The report does not present them as validated psychological scales.

The survey instrument was intentionally short to support participation. Each working domain in the emerging framework will need additional items and qualitative evidence before formal validation.

13. Conclusion

The Neuravox Human Agency Survey provides an initial empirical view of how everyday generative AI use relates to delegation, judgment, explanation and independent capability. The 76 respondents were experienced and frequent AI users. They reported moderate delegation alongside strong levels of conflict questioning and explanatory ownership. More than half perceived improvement in independent thinking, while about one quarter perceived some reduction.

The central finding is that frequency of AI use did not identify weaker human agency in this sample. Agency is better understood through the distribution of cognitive and decisional authority during AI use. Who starts the thinking, who structures the reasoning, who verifies the answer, who can explain the result, who retains the capability, and who makes the final decision are more informative questions.

These findings provide the empirical foundation for the next stage of the Neuravox Human Agency Framework. The framework will be developed as a set of distinct capabilities, tested through qualitative research, strengthened with better measurement, and validated in larger studies. The goal is to help people and institutions gain value from generative AI while keeping human judgment, skill and responsibility active.

References

- African Union. (2024). Continental Artificial Intelligence Strategy. African Union Commission. <https://www.au.int/en/documents/20240809/continental-artificial-intelligence-strategy>
- Buijsman, S., Carter, S. E., & Bermudez, J. P. (2025). Autonomy by Design: Preserving Human Autonomy in AI Decision-Support. arXiv:2506.23952. <https://arxiv.org/abs/2506.23952>
- Carolus, A., Koch, M., Straka, S., Latoschik, M. E., & Wienrich, C. (2023). MAILs: Meta AI Literacy Scale. arXiv:2302.09319. <https://arxiv.org/abs/2302.09319>
- Goddard, K., Roudsari, A., & Wyatt, J. C. (2012). Automation bias: a systematic review of frequency, effect mediators, and mitigators. *Journal of the American Medical Informatics Association*, 19(1), 121-127. <https://doi.org/10.1136/amiainl-2011-000089>
- Jin, Y., Martinez-Maldonado, R., Gasevic, D., & Yan, L. (2024). GLAT: The Generative AI Literacy Assessment Test. arXiv:2411.00283. <https://arxiv.org/abs/2411.00283>
- Lee, H. P. H., Sarkar, A., Tankelevitch, L., Drosos, I., Rintel, S., Banks, R., & Wilson, N. (2025). The Impact of Generative AI on Critical Thinking: Self-Reported Reductions in Cognitive Effort and Confidence Effects From a Survey of Knowledge Workers. *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*, 1-22. <https://doi.org/10.1145/3706598.3713778>
- OECD. (2024). OECD AI Principles. Organisation for Economic Co-operation and Development. <https://www.oecd.org/en/topics/ai-principles.html>
- Parasuraman, R., & Manzey, D. H. (2010). Complacency and Bias in Human Use of Automation: An Attentional Integration. *Human Factors*, 52(3), 381-410. <https://doi.org/10.1177/0018720810376055>
- Risko, E. F., & Gilbert, S. J. (2016). Cognitive Offloading. *Trends in Cognitive Sciences*, 20(9), 676-688. <https://doi.org/10.1016/j.tics.2016.07.002>
- Romeo, G., & Conti, D. (2026). Exploring automation bias in human-AI collaboration: a review and implications for explainable AI. *AI & Society*, 41, 259-278. <https://doi.org/10.1007/s00146-025-02422-7>
- Soto-Sanfiel, M. T., Angulo-Brunet, A., & Lutz, C. (2025). The scale of artificial intelligence literacy for all (SAIL4ALL): assessing knowledge of artificial intelligence in all adult populations. *Humanities and Social Sciences Communications*, 12, 1618. <https://doi.org/10.1057/s41599-025-05978-3>
- Sturgeon, B., Samuelson, D., Haimes, J., & Anthis, J. R. (2025). HumanAgencyBench: Scalable Evaluation of Human Agency Support in AI Assistants. arXiv:2509.08494. <https://arxiv.org/abs/2509.08494>
- UNESCO. (2021). Recommendation on the Ethics of Artificial Intelligence. United Nations Educational, Scientific and Cultural Organization. <https://www.unesco.org/en/legal-affairs/recommendation-ethics-artificial-intelligence>
- Zhu, Q., Li, X., Dong, Y., Chang, P., & Fan, M. (2026). Not all cognitive offloading is equal: distinguishing dependent and autonomous offloading to generative AI. *Frontiers in Psychology*, 17, 1878629. <https://doi.org/10.3389/fpsyg.2026.1878629>

Appendix A. Survey measures and coding

The seven behavioural items used the same response scale: Never, Rarely, Sometimes, Often and Almost always. The analysis coded these categories from 1 to 5.

Measure	Working dimension	Interpretation
Asks AI before self	Delegation	Higher score means more frequent use of AI before independent thinking.
Uses without checking	Delegation	Higher score means more frequent use of AI output without verification.
AI structures ideas	Delegation	Higher score means AI more often determines the structure of ideas or approach.
Delegates previous decisions	Delegation	Higher score means more frequent transfer of decisions previously made personally.
Questions conflicting AI	Retained agency	Higher score means more frequent challenge when AI conflicts with personal knowledge or judgment.
Can explain final work	Retained agency	Higher score means stronger reported ability to explain AI assisted work independently.
Confident without AI	Retained agency	Higher score means stronger confidence in completing similar tasks without AI.

Perceived change was coded from Reduced a lot through Improved a lot for ordered analysis. Not sure remained a separate category and was excluded from analyses that required an ordered outcome.

Appendix B. Selected statistical detail

Analysis	Result	Interpretation
Delegation score	Mean 2.71, SD 0.69	Four item exploratory summary
Retention score	Mean 4.02, SD 0.61	Three item exploratory summary
Overall candidate alpha	0.575	Composite not retained
Delegation alpha	0.641	Exploratory summary
Retention alpha	0.410	Interpret with individual items
Frequency vs Delegation	rho 0.046, p 0.692	No clear association
Frequency vs Retention	rho 0.247, p 0.031	Modest positive association
Frequency vs perceived change	rho 0.022, p 0.855	No clear association
Delegation vs perceived change	rho 0.010, p 0.934	No clear association
Retention vs perceived change	rho -0.257, p 0.031	Exploratory negative association

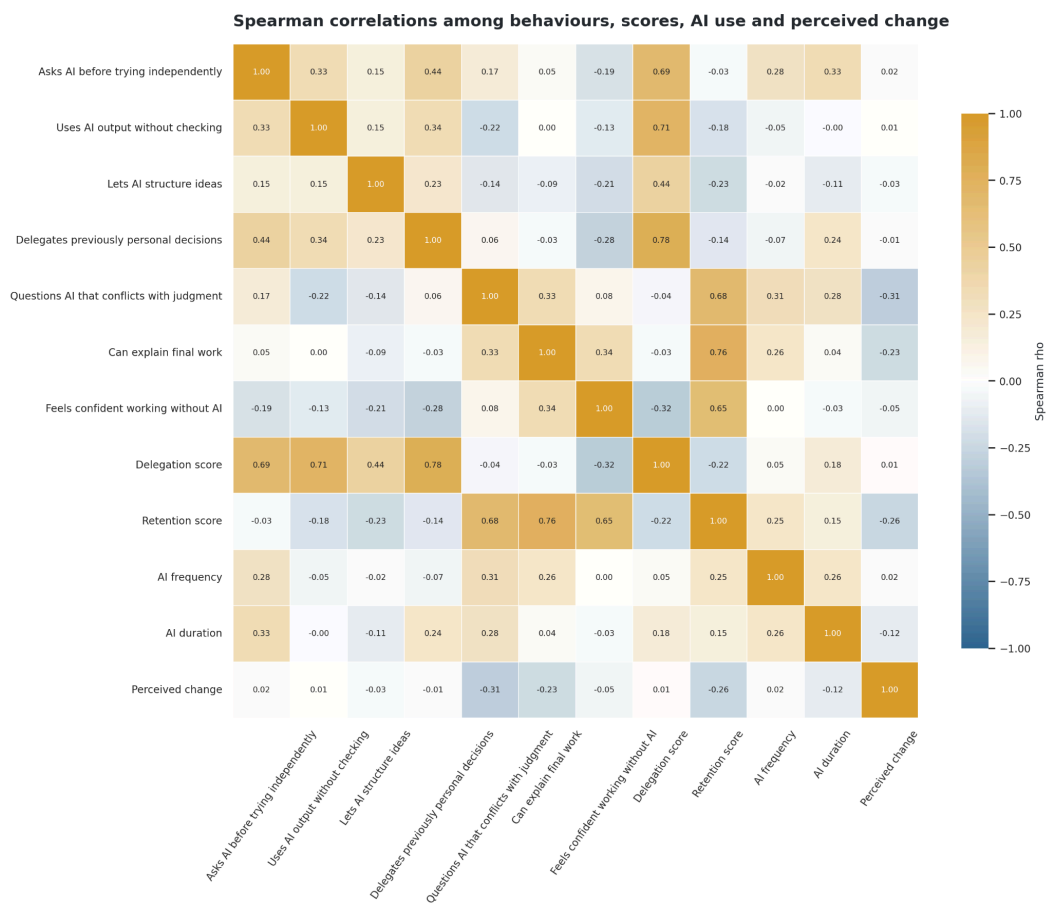


Figure A1. Spearman correlation heatmap for agency behaviours, exploratory scores, AI use frequency, AI use duration and perceived change.

The full reproducible analysis includes bootstrap confidence intervals, multiple testing correction, grouped comparisons, use case comparisons and a restrained ordinal regression. These analyses informed the report narrative and are available in the underlying project workspace.

END