

Aporia Intelligence (ApoQ): Theorizing Human Cognitive Advantage in AI-Augmented Organizations

Rahul Kale
Merieux Nutrisciences

Abstract - The proliferation of AI systems capable of generating analysis and judgment on demand is compressing the cognitive space in which human inquiry develops. This paper argues that the compression itself is the organizational problem worth examining, and introduces Aporia Intelligence (ApoQ) as a novel construct to name what is being lost. ApoQ is defined as the practiced, agentic, and temporally structured capacity to dwell productively in unresolved uncertainty in service of superior future judgment. Drawing on Socratic epistemology, the creativity and incubation literature, and the philosophy of agency, we develop ApoQ as a multidimensional construct with three facets, namely Recognition, Restraint, and Engagement. We distinguish it from tolerance of ambiguity (Frenkel-Brunswik, 1949) and epistemic humility (Krumrei-Mancuso and Rouse, 2016). We specify ApoQ's nomological network, including its antecedents, consequences, and boundary conditions, and offer a principled agentic argument for why ApoQ is structurally irreproducible by current AI systems. We connect this construct to AI hallucination, arguing that fluent but false AI outputs are what fluent-output systems produce in the absence of the agentic constraint that ApoQ instantiates in humans. Five testable propositions are developed across individual, team, and human-AI collaboration levels of analysis. The paper reframes the central question. Not what humans can do that AI cannot, but what cognitive disciplines humans are abandoning because AI makes them unnecessary.

Keywords: aporia intelligence; ApoQ; epistemic humility; AI hallucination; human-AI collaboration; organizational behavior; construct development; nomological network

INTRODUCTION

There is a question that management scholars have been circling for the better part of a decade, and it has become considerably more pressing in recent years. When artificial intelligence can do much of what knowledge workers do, what is it that knowledge workers are actually for?

The dominant answer, by now familiar, runs roughly as follows. AI is automating the cognitive outputs once associated with analytical intelligence, so the residual human advantage must lie in emotional intelligence. This includes empathy, relational trust, and the capacity to navigate social complexity in ways that machines cannot genuinely replicate (Goleman, 1995; Mayer et al., 2008). Advances in large language models have lent this argument a new urgency. Systems capable of passing the bar exam, synthesizing medical literature, and generating coherent strategic memos have arrived faster than most predicted (Bubeck et al., 2023, preprint; Singhal et al., 2023), and the premium on analytical capability has compressed accordingly.

We do not dispute this argument. However, we think it is incomplete in a way that has real organizational consequences.

What the IQ to EQ substitution narrative misses is a third cognitive capacity. It is neither about processing speed nor social fluency, but about something more fundamental: the capacity to stay in a question. To resist, deliberately and skillfully, the pull toward premature resolution. To recognize that the quality of eventual judgment depends on the quality of the inquiry that precedes it, and to act accordingly even when a ready answer is available.

We term this capacity Aporia Intelligence (ApoQ). The name is not accidental. Aporia, in the Platonic dialogues, names the productive state of genuine puzzlement that precedes authentic inquiry. It is not a failure state. It is the necessary threshold for real thinking to begin.

This paper makes five contributions. First, it introduces ApoQ as a novel, multidimensional construct developed according to the rigorous standards of Podsakoff, MacKenzie, and Podsakoff (2016), including formal specification of its essential attributes, dimensionality, and nomological network. Second, it carefully distinguishes ApoQ from two adjacent constructs, namely tolerance of ambiguity and epistemic humility. Third, it offers a principled agentive argument for why ApoQ is structurally irreproducible by current AI systems. Fourth, it connects ApoQ to the phenomenon of AI hallucination, offering an organizational lens that complements the established technical explanations of hallucination in the computer science literature. This connection in turn identifies a specific organizational role for high-ApoQ humans in AI-augmented workflows. Fifth, and perhaps most broadly, it reframes the central question driving human-AI research. The question is not what humans can do that AI cannot, but what cognitive disciplines humans are abandoning because AI makes them unnecessary, and at what organizational cost.

THEORETICAL BACKGROUND

The Compression of Cognitive Space by AI

Contemporary AI systems are designed, at a very deep level, to produce fluent and contextually coherent output in response to any input (Brown et al., 2020; Radford et al., 2019). They resolve uncertainty continuously and near-instantaneously. This is what we want from them for most tasks. However, embedding these systems in decision support, problem formulation, and knowledge synthesis workflows has a consequence that has received insufficient attention in the organizational behavior literature. It structurally abbreviates the cognitive experience of uncertainty, that is, the productive discomfort of not yet knowing, which is the input to ApoQ.

The organizational consequence is not merely one of efficiency. When uncertainty is resolved before it can be fully inhabited, the incubation phase of cognition, which the creativity literature identifies as a necessary precondition for genuine insight, is bypassed (Wallas, 1926; Csikszentmihalyi, 1990). As Daugherty and Wilson (2018) argue, the future of work lies not in human versus machine but in human plus machine. This framing assumes humans bring something distinctive to the collaboration. ApoQ names one critical element of what that distinctive contribution is, and why it is at risk of being systematically eroded.

Creativity, Incubation, and the Productive Role of Unresolved Uncertainty

The creativity literature has long recognized that not all phases of problem-solving benefit from active, convergent effort. Wallas (1926) identified incubation as a discrete phase of creative cognition during which subconscious processing continues in the absence of directed attention. What incubation does, we now understand with neuroscientific precision, involves the spontaneous formation of novel associative connections across conceptual domains that focused cognition tends to suppress (Beaty et al., 2016; Jung-Beeman et al., 2004). Csikszentmihalyi (1990) found that sustained engagement with a problem's unresolved aspects was a precondition for the generative states from which insight emerges.

ApoQ is the capacity to make this incubation intentional. It is not passive waiting, but rather the recognition that suspension of resolution serves the quality of eventual judgment, and the maintenance of productive engagement with the open question during that interval. This is a discipline, and like all disciplines, it can be developed or atrophied. In AI-saturated cognitive environments, we argue that the structural conditions for ApoQ are being systematically undermined. This is a theoretical claim we develop below, and that motivates the empirical research program outlined in the propositions.

Adjacent Constructs: Tolerance of Ambiguity and Epistemic Humility

Two constructs in the existing literature are close enough to require careful distinction before introducing ApoQ formally.

Tolerance of ambiguity (TA), introduced by Frenkel-Brunswik (1949) and operationalized by Budner (1962) and McLain (1993), describes an individual difference in how threatening a person finds ambiguous, uncertain, or novel situations. It is primarily a stable trait, shaped by early developmental experience and personality structure (Furnham and Ribchester, 1995).

Epistemic humility (EH), as operationalized by Krumrei-Mancuso and Rouse (2016), is the metacognitive awareness that one's own knowledge and beliefs may be incomplete or wrong. It is a cognitive orientation toward one's own epistemic state.

ApoQ is neither of these. Tolerance of ambiguity describes a disposition toward uncertainty, that is, how you feel about it. Epistemic humility describes a metacognitive stance, that is, what you know about the limits of your knowledge. ApoQ describes a

practice, that is, what you do inside uncertainty, deliberately and in service of a purpose. These distinctions are formalized in Table 1.

Table 1 - Comparison of Adjacent Constructs

Dimension	Tolerance of Ambiguity	Epistemic Humility	ApoQ
Type	Stable trait	Metacognitive awareness	Multidimensional practice
Core concern	Comfort with unclear situations	Knowing limits of one's knowledge	Dwelling productively in uncertainty
Orientation	Dispositional	Cognitive	Agentic, volitional
AI replicability	Partially simulable	Partially simulable	Structurally absent
Development	Trait; partial plasticity	Trainable via reflection	Explicitly trainable; skill-based

THE APOQ CONSTRUCT

Formal Definition

Aporia Intelligence (ApoQ) is the practiced, agentic, and temporally structured capacity to dwell productively in a state of unresolved cognitive uncertainty, deliberately suspending premature closure, in service of superior future judgment.

The term *aporia* derives from the Platonic dialogues, where it names the productive crisis of recognizing that one does not know what one thought one knew. In the *Meno*, Socrates induces *aporia* not to humiliate his interlocutor but because genuine inquiry cannot begin until the question is genuinely open. *Aporia* is the doorway, not the destination (Plato, trans. Grube, 1981). ApoQ operationalizes this insight as a trainable cognitive practice rather than a philosophical state.

Essential Attributes of ApoQ

Following Podsakoff et al. (2016), we specify the essential attributes that are common and unique to ApoQ as a construct. Four attributes are necessary and sufficient.

1. Agentic. ApoQ involves a deliberate choice to remain in uncertainty, not a passive tolerance of it.
2. Temporally structured. ApoQ is not an indefinite state. It has a beginning (recognition) and an end (resolution, when the agent judges that inquiry has been sufficiently served).
3. Instrumentally motivated. ApoQ is exercised in service of better eventual judgment, not as intellectual asceticism.
4. Trainable. ApoQ is a practice rather than a fixed trait. It can be developed through deliberate intervention and atrophied through disuse.

Dimensionality: A Multidimensional Construct

Podsakoff et al. (2016) require explicit specification of whether a construct is unidimensional or multidimensional. We propose that ApoQ is a multidimensional construct comprising three related but conceptually distinct facets.

Recognition

The cognitive ability to identify when a question is genuinely unresolved in a way that carries meaningful consequences, and when the costs of premature resolution outweigh the benefits of speed.

Restraint

The volitional capacity to withhold resolution when it is available, resisting the organizational and technological pressures that favor fast answers. This facet is the agentic core of ApoQ.

Engagement

The capacity to maintain productive inquiry with an open question during the interval of deferred resolution, rather than simply waiting passively.

These three facets are related. Individuals with high ApoQ tend to score highly on all three. However, they are conceptually distinguishable and likely to show differential relationships with antecedents and outcomes. For example, Recognition may be more strongly predicted by need for cognition and metacognitive awareness, while Restraint may be more strongly related to self-regulation capacity (Baumeister and Vohs, 2004) and tolerance of ambiguity. We treat ApoQ as a higher-order construct in which the three facets are related reflective indicators of the overall construct.

Why AI Cannot Exhibit ApoQ

The argument that ApoQ is structurally irreproducible by current AI systems is not architectural. It does not rest on technical limitations that better models will close. The argument is agentive, grounded in Bratman's (1987) theory of planning agency and Dennett's (1987) intentional stance.

ApoQ's three facets each require something AI systems do not possess. Recognition requires caring about the quality of one's eventual judgment in a way that motivates behavioral change. Restraint requires choosing not to produce output when output is available, because of an agent's stake in the quality of what comes next. Engagement requires maintaining a purposive relationship with an open question over time. Current AI systems are designed to produce output in response to input. The suppression of output in favor of productive waiting is not part of their behavioral repertoire, not because they lack the technical capacity to delay response, but because they lack the stakes. Nothing is invested in the quality of their eventual judgment.

This argument does not require claims about consciousness or general intelligence. It requires only that motivated restraint, that is, choosing not to resolve because one cares about the quality of eventual resolution, is a form of agency that current and near-future AI systems do not exhibit.

This argument connects to, but is distinct from, earlier philosophical critiques of AI cognition. Searle's (1980) Chinese Room argument challenged whether symbol manipulation constitutes genuine understanding. Dreyfus (1972) questioned whether human expertise could be captured in formal rules. Our claim is narrower and more specific. We do not argue that AI lacks understanding in some deep sense, but that it lacks the motivational structure (stakes, investment in outcome quality, behavioral consequences of judgment failure) that makes ApoQ's volitional restraint possible.

AI Hallucination in the Absence of Agentive Constraint

The theoretical framework developed above has a concrete and immediately consequential manifestation in one of the most widely discussed limitations of contemporary AI systems, namely hallucination. Hallucination is the tendency of large language models to generate fluent, confident, and factually incorrect outputs (Ji et al., 2023; Zhang et al., 2023). Before offering our organizational lens on this phenomenon, we acknowledge the well-established technical explanations. Computer science research attributes hallucination to training data limitations, likelihood-based decoding procedures that select fluent tokens without factual grounding, and the absence of retrieval mechanisms that would tie outputs to verified sources (Ji et al., 2023; Kadavath et al., 2022). These technical accounts are correct and consequential. Our contribution is not to displace them but to add a complementary organizational lens that clarifies why hallucination persists as an organizational problem even as technical mitigations improve, and what role humans play in the mitigation.

We wish to be careful about a category question. ApoQ is defined as an agentive construct that requires stakes, investment in outcome quality, and behavioral consequences of judgment failure. It would therefore be a category error to say AI systems lack ApoQ in the same sense that a person lacks it, since AI systems lack the entire agentive substrate in which ApoQ operates. The more precise claim is this. Hallucination is what fluent-output systems produce in the absence of any agentive constraint on output generation. ApoQ is the human form of such constraint. Technical calibration mechanisms are a different form. The two are not equivalent, and understanding the difference has practical consequences for how organizations should design human-AI workflows.

Consider the phenomenon carefully. When an AI system produces a hallucinated output, it is not deceiving the user in any morally meaningful sense. It has no intent. What it is doing is producing the most probabilistically fluent continuation of the input,

unconstrained by any mechanism that would cause it to hold the continuation open when it exceeds its epistemic warrant. Mapping the absence of ApoQ-like constraint onto its three facets makes the claim structurally precise.

1. Absence of Recognition-analogous mechanism. The system does not identify when a question is genuinely difficult, when the available information does not support a confident answer, or when the epistemic status of its response should be qualified. It treats questions of vastly different reliability with equivalent fluency.
2. Absence of Restraint-analogous mechanism. The system does not withhold a response when withholding is warranted. Given input, it produces output. It does not possess the volitional architecture that would produce 'I do not know,' 'this exceeds what I can defensibly claim,' or 'let me hold this open until I have thought further.'
3. Absence of Engagement-analogous mechanism. There is no interval of productive inquiry between question and answer. The system does not sit with the question, explore its edges, or allow associative processing to continue. There is only input and output.

This framing clarifies ongoing efforts to reduce hallucination through improved calibration, retrieval-augmented generation, and confidence estimation (Kadavath et al., 2022). These are valuable technical interventions that partially substitute for the missing agentive constraint. Our framework predicts that they will produce meaningful but bounded improvements. Calibration is a statistical property of outputs. ApoQ is an agentive practice grounded in stakes and consequences. Even a perfectly calibrated AI system, in the sense of accurately reporting its uncertainty, would still lack the motivational structure to hold an open question when doing so would serve the quality of the user's eventual judgment. This yields a falsifiable prediction. If our framework is correct, hallucination rates and, more importantly, the acceptance rates of hallucinated outputs by users, will remain non-negligible in consequential decision contexts even as technical calibration approaches optimal levels, unless organizational workflow design activates human ApoQ as a complementary constraint. If technical calibration alone eliminated hallucination acceptance in consequential decisions, our framework's contribution would be substantially weakened.

The organizational implication follows directly. High-ApoQ humans working alongside AI systems serve a specific epistemic function that low-ApoQ humans cannot. Within their domains of competence, they can recognize when an AI output warrants suspicion, restrain themselves from accepting fluent but underconstrained content, and engage productively with the uncertainty rather than defaulting to the AI's fluent confidence. We are careful to specify within their domains of competence, because ApoQ moderates but does not substitute for domain expertise. A high-ApoQ humanities professor evaluating a hallucinated code snippet lacks the substantive knowledge to detect the hallucination regardless of their cognitive discipline. ApoQ therefore functions as an epistemic amplifier within one's competence area, not a general defense against all fluent but false content.

Nomological Network

A complete construct definition specifies what causes a construct, what it causes, and what moderates these relationships (Podsakoff et al., 2016). We specify ApoQ's nomological network as follows.

Antecedents

Individual differences predicting higher ApoQ include need for cognition (Cacioppo and Petty, 1982), intellectual humility, and mindfulness. Developmentally, exposure to philosophical inquiry, Socratic dialogue pedagogy, and reflective practice traditions should predict higher ApoQ. Situationally, environments that reward fast answers and penalize visible uncertainty are expected to suppress ApoQ expression.

Consequences

ApoQ is predicted to produce superior decision quality under conditions of genuine uncertainty (Knightian uncertainty), higher creative output on novel problems, greater resistance to overconfident and hallucinated AI recommendations within one's domain of competence, and more effective leadership in high-ambiguity contexts.

Boundary Conditions

ApoQ's positive effects are strongest under conditions of high task novelty, when AI outputs are overconfident or prone to hallucination, when the decision falls within the individual's domain of competence, and when organizational stakes are high. Under low uncertainty or high time pressure with real costs to delay, ApoQ's value diminishes.

Table 2 - ApoQ Nomological Network

Category	Variable	Direction and Nature of Relationship
Antecedents	Need for cognition	Positive: intrinsic motivation to engage in effortful thinking predicts higher ApoQ
	Intellectual humility	Positive: openness to not-knowing supports uncertainty-dwelling
	Mindfulness	Positive: present-moment awareness facilitates productive engagement
	AI tool availability	Negative: high AI availability increases pressure toward premature resolution
Consequences	Decision quality under uncertainty	Positive: ApoQ predicts superior judgment
	Creative insight	Positive: ApoQ facilitates incubation
	Hallucination resistance (within competence)	Positive: high ApoQ reduces acceptance of fluent but false AI outputs within one's domain of expertise
	Leadership effectiveness	Positive in high-ambiguity contexts
Boundary conditions	Task novelty	Amplifier: effects strongest on novel tasks
	Time pressure	Attenuator: high time pressure reduces ApoQ value
	Domain expertise	Necessary co-moderator: ApoQ operates as epistemic amplifier only within competence
	AI confidence and hallucination risk	Amplifier: effects strongest when AI outputs are overconfident or prone to hallucination

The Three-Construct Model: IQ, EQ, and ApoQ

The three constructs represent a changing distribution of scarcity in AI-augmented organizations. The economic outputs associated with high IQ (rapid analysis, pattern recognition, structured synthesis) are now largely automatable by AI systems, compressing the premium on raw analytical horsepower regardless of what IQ itself measures neurologically. EQ-related capacities are partially simulable. ApoQ is structurally resistant to automation for the agentive reasons specified above.

Anchoring this in resource-based view theory (Barney, 1991), ApoQ qualifies as a rare, inimitable, and non-substitutable cognitive resource whose strategic value increases precisely as AI commoditizes adjacent capabilities. From an attention-based view (Ocasio, 1997), AI-saturated environments systematically make fast answers more salient than productive uncertainty. This dynamic makes ApoQ both rarer and more valuable as AI deployment deepens.

RESEARCH PROPOSITIONS

Five propositions follow from the theoretical framework and the nomological network specified above.

Individual Level

Proposition 1: *ApoQ will be positively associated with decision quality under conditions of genuine uncertainty, over and above the contributions of IQ, EQ, tolerance of ambiguity, and need for cognition.*

Proposition 2: *Structured ApoQ-development interventions incorporating Socratic dialogue, philosophical inquiry, and deliberate incubation protocols will produce significant and durable increases in ApoQ scores relative to active control conditions.*

Team and Leadership Level

Proposition 3: *Leaders who publicly acknowledge unresolved uncertainty with authenticity and frequency will be rated as more trustworthy and more effective by direct reports, particularly in novel and high-ambiguity task environments, and this relationship will be mediated by follower perceptions of leader ApoQ. We propose this mediation because followers interpret deliberate uncertainty acknowledgment as evidence that their leader possesses the judgment quality associated with productive uncertainty-dwelling. This constitutes a form of cognitive credibility signaling that is distinct from impression management.*

Proposition 4: *Team-level ApoQ will moderate the relationship between AI tool availability and team decision quality, such that high-ApoQ teams will outperform low-ApoQ teams under conditions of high AI availability.*

Human-AI Collaboration and Hallucination Resistance

Proposition 5: *Within their domains of competence, high-ApoQ individuals will show lower acceptance rates of hallucinated or overconfident AI outputs on genuinely uncertain tasks, and will achieve superior objective accuracy relative to low-ApoQ individuals under equivalent conditions. Domain expertise is a necessary co-moderator. ApoQ operates as an epistemic amplifier within one's competence area, not a general defense across all subject matters. This proposition is directly testable using laboratory paradigms in which participants complete tasks with AI assistance that has been experimentally seeded with fluent but incorrect outputs, with participants sampled from populations in which the seeded errors fall within their substantive expertise. High-ApoQ participants are predicted to detect and reject these outputs at higher rates than low-ApoQ participants matched on domain expertise.*

THEORETICAL AND PRACTICAL IMPLICATIONS

Theoretical Contributions

This paper makes five theoretical contributions. The first is the introduction of ApoQ as a novel, rigorously specified construct. The second is the principled agentic argument for ApoQ's structural irreproducibility by AI systems. The third is the scarcity-based three-construct model organizing IQ, EQ, and ApoQ around a dimension of AI replicability. The fourth is the connection between ApoQ and AI hallucination, offering an organizational lens that complements established technical accounts, generating a falsifiable prediction about the limits of technical calibration alone, and identifying an actionable organizational role for high-ApoQ humans in AI-augmented workflows, formalized in Proposition 5. The fifth and perhaps most significant contribution is a reframing of the central question. The question shifts from 'what can humans do that AI cannot?' to 'what cognitive disciplines are humans abandoning because AI makes them unnecessary, and at what organizational cost?'

Practical Implications

For leadership development, ApoQ suggests that cultivating the capacity to publicly inhabit unresolved uncertainty is a cognitive discipline with measurable consequences for judgment quality. Executive coaching curricula incorporating Socratic dialogue and structured incubation are investments in a form of capability that is becoming strategically scarcer.

For team design and AI deployment, the moderating role of team-level ApoQ suggests a previously unrecognized composition consideration. Organizations deploying AI decision-support tools may need to actively assess and develop collective ApoQ as a counterweight to the premature-closure and hallucination-acceptance risks embedded in high-AI-availability environments.

For human-AI workflow design, the connection between ApoQ and hallucination acceptance suggests specific interventions. Organizations can (a) design workflows that prompt explicit ApoQ engagement before accepting AI outputs on consequential decisions within one's domain of competence, (b) develop training that helps employees recognize AI output patterns most susceptible to hallucination in their specific technical domains, and (c) treat ApoQ development as a specific competency to be assessed and cultivated alongside AI fluency and domain expertise. High AI fluency without ApoQ produces employees who

accept fluent but false outputs quickly. AI fluency, ApoQ, and domain expertise together produce employees who can extract AI's value while catching its errors.

LIMITATIONS AND FUTURE DIRECTIONS

This paper is conceptual. The propositions advanced here are theoretically grounded but empirically untested. Future psychometric work should develop and validate a formal measurement instrument for the three-facet structure specified in Section 3, following established scale-development methodology (Hinkin, 1998); we intentionally do not undertake that separate empirical project here, and we treat it as an important and substantial direction for future research rather than a claim this paper is positioned to make.

A second limitation concerns cultural specificity. The Socratic tradition on which the construct rests is Western in origin. Hofstede's (1980) uncertainty avoidance dimension provides an anchor for cross-cultural validity work.

A third limitation is the boundary condition specification, which remains theoretical. The conditions under which ApoQ imposes costs, such as high time pressure with real penalties for delay, deserve explicit empirical attention.

A fourth limitation concerns the heterogeneity of AI hallucination. Hallucinations vary in type, including factual errors, fabricated citations, invented events, coherent but wrong reasoning chains, and confidently stated opinions on subjective questions. Our framework treats hallucination categorically. Whether ApoQ predicts resistance equally across all types, or differentially, is an empirical question we do not resolve here.

Future directions include the psychometric scale development work noted above, experimental ApoQ training interventions, field studies in AI-augmented organizations across industries, laboratory studies testing Proposition 5 using experimentally-seeded hallucinated AI outputs across different hallucination types, and cross-cultural validation.

CONCLUSION

We began with a question that has become unavoidable in organizational research. As AI systems perform more of what knowledge workers do, what remains distinctively and strategically human about organizational cognition? The standard answer, emotional intelligence, is important but incomplete.

This paper has introduced Aporia Intelligence (ApoQ) as a novel construct that fills a gap the EQ narrative does not address. ApoQ is a rigorously specified, multidimensional, trainable cognitive practice with identifiable antecedents, measurable consequences, and a principled account of why it cannot be replicated by the AI systems that are compressing the cognitive space in which it operates. We have shown that ApoQ offers explanatory purchase on one of the most consequential contemporary AI phenomena, namely hallucination, and that this connection identifies a specific organizational role for high-ApoQ humans in AI-augmented workflows within their domains of competence.

The reframe we offer is this. The question is not what humans can do that AI cannot. It is what cognitive disciplines humans are abandoning because AI makes them unnecessary. In a world where answers are available on demand, the capacity to resist them, to hold the question open, to notice the fluent but false, to stay in productive uncertainty, is what becomes scarce. ApoQ names that capacity.

REFERENCES

- [1] Barney, J. B. (1991). Firm resources and sustained competitive advantage. *Journal of Management*, 17(1), 99-120.
- [2] Baumeister, R. F., and Vohs, K. D. (2004). *Handbook of self-regulation: Research, theory, and applications*. Guilford Press.
- [3] Beaty, R. E., Benedek, M., Silvia, P. J., and Schacter, D. L. (2016). Creative cognition and brain network dynamics. *Trends in Cognitive Sciences*, 20(2), 87-95.
- [4] Bratman, M. E. (1987). *Intention, plans, and practical reason*. Harvard University Press.
- [5] Brown, T., Mann, B., Ryder, N., Subbiah, M., Kaplan, J. D., Dhariwal, P., and Amodei, D. (2020). Language models are few-shot learners. *Advances in Neural Information Processing Systems*, 33, 1877-1901.
- [6] Bubeck, S., Chandrasekaran, V., Eldan, R., Gehrke, J., Horvitz, E., Kamar, E., and Zhang, Y. (2023). Sparks of artificial general intelligence: Early experiments with GPT-4. *arXiv preprint arXiv:2303.12712*. Manuscript under review.
- [7] Budner, S. (1962). Intolerance of ambiguity as a personality variable. *Journal of Personality*, 30(1), 29-50.
- [8] Cacioppo, J. T., and Petty, R. E. (1982). The need for cognition. *Journal of Personality and Social Psychology*, 42(1), 116-131.
- [9] Csikszentmihalyi, M. (1990). *Flow: The psychology of optimal experience*. Harper and Row.
- [10] Daugherty, P. R., and Wilson, H. J. (2018). *Human plus machine: Reimagining work in the age of AI*. Harvard Business Review Press.

- [11] Dennett, D. C. (1987). *The intentional stance*. MIT Press.
- [12] Dreyfus, H. L. (1972). *What computers cannot do: A critique of artificial reason*. Harper and Row.
- [13] Frenkel-Brunswick, E. (1949). Intolerance of ambiguity as an emotional and perceptual personality variable. *Journal of Personality*, 18(1), 108-143.
- [14] Furnham, A., and Ribchester, T. (1995). Tolerance of ambiguity: A review of the concept, its measurement and applications. *Current Psychology*, 14(3), 179-199.
- [15] Goleman, D. (1995). *Emotional intelligence: Why it can matter more than IQ*. Bantam Books.
- [16] Hinkin, T. R. (1998). A brief tutorial on the development of measures for use in survey questionnaires. *Organizational Research Methods*, 1(1), 104-121.
- [17] Hofstede, G. (1980). *Culture's consequences: International differences in work-related values*. Sage.
- [18] Ji, Z., Lee, N., Frieske, R., Yu, T., Su, D., Xu, Y., Ishii, E., Bang, Y. J., Madotto, A., and Fung, P. (2023). Survey of hallucination in natural language generation. *ACM Computing Surveys*, 55(12), 1-38.
- [19] Jung-Beeman, M., Bowden, E. M., Haberman, J., Frymiare, J. L., Arambel-Liu, S., Greenblatt, R., and Kounios, J. (2004). Neural activity when people solve verbal problems with insight. *PLOS Biology*, 2(4), e97.
- [20] Kadavath, S., Conerly, T., Askill, A., Henighan, T., Drain, D., Perez, E., Schiefer, N., et al. (2022). Language models (mostly) know what they know. *arXiv preprint arXiv:2207.05221*.
- [21] Krumrei-Mancuso, E. J., and Rouse, S. V. (2016). The development and validation of the Comprehensive Intellectual Humility Scale. *Journal of Personality Assessment*, 98(2), 209-221.
- [22] Mayer, J. D., Roberts, R. D., and Barsade, S. G. (2008). Human abilities: Emotional intelligence. *Annual Review of Psychology*, 59, 507-536.
- [23] McLain, D. L. (1993). The MSTAT-I: A new measure of an individual's tolerance for ambiguity. *Educational and Psychological Measurement*, 53(1), 183-189.
- [24] Ocasio, W. (1997). Towards an attention-based view of the firm. *Strategic Management Journal*, 18(S1), 187-206.
- [25] Plato. (1981). *Meno* (G. M. A. Grube, Trans.). Hackett Publishing. (Original work composed c. 385 BCE)
- [26] Podsakoff, P. M., MacKenzie, S. B., and Podsakoff, N. P. (2016). Recommendations for creating better concept definitions in the organizational, behavioral, and social sciences. *Organizational Research Methods*, 19(2), 159-203.
- [27] Radford, A., Wu, J., Child, R., Luan, D., Amodei, D., and Sutskever, I. (2019). Language models are unsupervised multitask learners. *OpenAI Blog*, 1(8), 9.
- [28] Searle, J. R. (1980). Minds, brains, and programs. *Behavioral and Brain Sciences*, 3(3), 417-424.
- [29] Singhal, K., Azizi, S., Tu, T., Mahdavi, S. S., Wei, J., Chung, H. W., and Natarajan, V. (2023). Large language models encode clinical knowledge. *Nature*, 620(7972), 172-180.
- [30] Wallas, G. (1926). *The art of thought*. Harcourt Brace.
- [31] Zhang, Y., Li, Y., Cui, L., Cai, D., Liu, L., Fu, T., Huang, X., Zhao, E., Zhang, Y., Chen, Y., Wang, L., Luu, A. T., Bi, W., Shi, F., and Shi, S. (2023). Siren's song in the AI ocean: A survey on hallucination in large language models. *arXiv preprint arXiv:2309.01219*.