

The Cognitive Sovereignty Project

Competent Output Is Getting Cheap. Independent Judgment Is Not.

AI democratizes output faster than it democratizes judgment. The same tool amplifies a formed mind, extends a partly formed one, and substitutes for one that was never built — and from the finished artifact alone, the three are hard to tell apart. This paper proposes prior formation as the candidate mechanism, engages the strongest published evidence against that position by name and number, and publishes eight conditions under which its own framework should be abandoned.

Yvonne G. Hall, Founder & Chief Executive Officer, EQGenix Inc. · Canyon Lake, California · August 2026 · Research Edition v3.3

Contents

A Note on Evidence

A Note on Method

I. The Premium Is Moving

II. Five Findings, Five Methods, One Unresolved Boundary

III. Three Mechanisms of Erosion

IV. The Smoothness Trap™

V. The Relocation of Authority™

VI. Amplifier, Collaborator, Substitute

- VII. The Strongest Counterargument
- VIII. Cognitive Sovereignty™ Defined
- IX. The Four Disciplines of Rebuilding
- X. The Foundation Beneath the Framework
- XI. What This Means for Institutions
- XII. The Rising Premium on Independent Judgment
- XIII. What Would Have to Be True
- XIV. Governance Boundaries
- XV. What This Paper Does Not Claim
- XVI. Research Protocol
- References

The Central Proposition

AI can democratize competent output without democratizing judgment. Formation influences whether the tool functions as an amplifier, a collaborator, or a substitute.

Executive Brief

Cognitive Sovereignty™ is the developed capacity to originate, hold, test, revise, and defend one's own judgment while using powerful external systems. EQGenix argues that AI does not produce one uniform human outcome. It can amplify developed judgment, widen access to competent execution, or substitute for capacities that were never formed. Which outcome occurs is not solely a property of the technology. The operator and the sequence appear to matter substantially.

The evidence base is consequential but still emerging, and this paper treats it that way. A peer-reviewed CHI survey links

confidence in AI to reduced self-reported critical engagement. An observational UK study associates frequent AI use, cognitive offloading, and lower critical-thinking scores. A preliminary EEG preprint reports different engagement patterns by tool condition and by sequence. Established memory research shows that expected access changes what people encode, and a peer-reviewed navigation study shows a delegated spatial capacity thinning over time. A peer-reviewed preregistered field experiment with 758 consultants shows AI lifting performance inside its capability frontier and degrading it outside — a boundary elite professionals could not reliably perceive. No single study demonstrates broad cognitive decline. Together they justify deliberate measurement and deliberate formation.

The institutional decision: automate execution without accidentally automating the experiences that develop judgment. Protect live reasoning, evidence tracing, adversarial testing, reflection, and accountable decision-making wherever those activities are load-bearing.

The EQGenix response: four disciplines — Position Before Prompt, Sequenced Difficulty, Adversarial Intake, and Sourced Conviction — supported by five observable markers and an assessment architecture under development for youth, families, employers, educators, public safety, and veteran transition.

Abstract

The prevailing anxiety about artificial intelligence is that it will make human beings unnecessary. The evidence points somewhere more uncomfortable. AI may compress differences in visible output while widening differences in the judgment

beneath it. Across five independent lines of research published between 2011 and 2026, a consequential pattern surfaces: powerful cognitive tools do not distribute their benefits evenly. They can amplify minds formed before the tool arrived while substituting for capacities that remain underdeveloped. The cited evidence does not identify a single differentiating variable. This paper proposes prior formation as a candidate mechanism requiring direct measurement, and Section XIII states what would show it is the wrong one.

This paper proposes that persistent cognitive offloading may make independently governed judgment scarcer relative to demand in the labor market, the civic square, and the household. Neither the direction nor the magnitude of that change has been established, and the proposition is advanced as a hypothesis rather than a finding. It identifies three mechanisms driving its erosion, names two failure patterns that precede collapse, and proposes four disciplines by which the capacity can be deliberately rebuilt. It engages, directly and by name, the strongest published evidence against its own thesis. The argument is not that AI is dangerous. The argument is that AI is extraordinary, and that extraordinary tools require sovereign operators.

A Note on Evidence

This is a founder-led research and perspective paper, not peer-reviewed institutional research, and it should be read as such. Every citation has been verified against the published version of record. Study limitations — preprint status, sample size, design, journal tier, published corrections and methodological

commentary — are disclosed inline at the point each finding is used rather than buried at the end. Where the evidence supports only an association, the language says association.

This edition makes that discipline auditable rather than merely asserted. Every substantive proposition carries one of four evidence classes, marked in the text at the point of use.

ESTABLISHED EVIDENCE means peer-reviewed findings supported by mature literature. EMERGING EVIDENCE means peer-reviewed but recent, correlational, narrow, or context-specific. PRELIMINARY EVIDENCE means preprints, small samples, or exploratory findings awaiting replication. EQGENIX ARGUMENT means interpretation or prediction advanced by this paper beyond what the cited evidence establishes.

Where a proposition carries the EQGENIX ARGUMENT mark, the reader is being told plainly that the sentence advances a position the cited evidence does not establish. Section VII engages the strongest published evidence against the thesis by name and number.

A Note on Method

This edition was produced under adversarial review. The preceding version was submitted to two independent reviewers applying different standards — one evaluating evidentiary defensibility, one evaluating institutional and commercial utility. The reviews disagreed materially, including on whether the assessment architecture should be expanded or withheld. Disagreements were adjudicated against the stated evidence

standard above rather than resolved by accepting both. Separately, and after review, every citation was re-verified against the published version of record; that pass caught two errors neither review identified. Reviewers advise; the evidence standard decides.

Earlier editions referred readers to a companion research agenda, available on request, for the conditions under which the EQGenix position would be weakened. Those conditions are now published in Section XIII of this paper, where they can be held against the firm without anyone having to ask for them.

I. The Premium Is Moving

EQGENIX ARGUMENT

Technological history repeatedly shows a shift in scarcity: when a capability becomes cheaper to execute, value tends to migrate toward the judgment required to direct, evaluate, or distinguish that execution. The printing press collapsed the cost of copying text and raised the premium on having something worth copying. The calculator collapsed the cost of arithmetic and raised the premium on knowing which calculation to run. The search engine collapsed the cost of retrieval and raised the premium on judgment about what was worth retrieving.

Artificial intelligence has sharply reduced the time and cost of producing competent drafts, summaries, code, analyses, structures, and translations. In many bounded tasks, work that once required professional time can now be generated in seconds. This is not a threat to be managed. It is among the most significant expansions of individual capability in the

modern era, and this paper concedes nothing on that point.

But the premium has moved. As competent output becomes ambient and inexpensive, competent output becomes less of what anyone pays for. What remains scarce is the capacity that sits upstream of the output: the ability to determine what question is worth asking, to recognize when a fluent answer is a wrong answer, to hold a position under pressure from a system that sounds more confident than you do, and to originate a direction that no model was trained to produce.

That capacity is not intelligence, which is raw processing. It is not education, which is credentialed exposure. It is sovereignty — the developed, defensible ownership of one's own thinking. This paper proposes that it is being spent down faster than it is being built. That is a hypothesis about a trend nobody has measured, and Section XIII states what would show it is wrong.

II. Five Findings, Five Methods, One Unresolved Boundary

Five findings, drawn from independent research teams using different methods on different populations, bear on a common question. They do not all point the same way — the productivity experiments in Section VII show substantial democratizing benefits — and what they jointly suggest is conditionality rather than uniform erosion. They are presented in the order they compound, each with its limitations stated where the finding is used.

Finding 1 — Trust in the tool predicts the abdication

EMERGING EVIDENCE

Researchers at Microsoft Research and Carnegie Mellon University surveyed 319 knowledge workers who use generative AI at least weekly, collecting 936 first-hand accounts of AI-assisted work tasks. The central result is not that AI reduces thinking. It is what predicted whether critical engagement was reported at all.

Two variables predicted critical engagement, and they pulled in opposite directions. Higher confidence in the AI was associated with less critical thinking. Higher confidence in one's own domain expertise was associated with more — despite the additional effort it cost. Workers reporting less confidence in their own expertise were less likely to report interrogating what the system produced.

Limitation disclosed at point of claim: this study is correlational and self-reported. Participants described their own cognitive effort rather than being observed or tested, and self-report on one's own thinking is not a strong instrument. The finding cannot establish that AI confidence causes disengagement. What it establishes is an association pointing in a specific direction — the tool alone did not determine the reported level of critical engagement, and confidence in the tool and confidence in one's own expertise moved behavior in opposite directions.

Finding 2 — Cognitive offloading may have a neural signature

PRELIMINARY EVIDENCE

A 2025 MIT Media Lab preprint examined the question using electroencephalography. Fifty-four participants wrote essays across three sessions under one of three conditions: using a

large language model, using a search engine, or using nothing but their own minds. In a fourth session the researchers reversed the assignments.

Brain connectivity separated along the lines you would expect and one line you would not. The unassisted writers showed the strongest and most distributed neural networks. Search engine users showed moderate engagement. LLM users showed the weakest connectivity of the three. Reported cognitive engagement scaled down in proportion to how much of the task was handed outward. LLM users also reported the lowest sense of ownership over their own essays and struggled to quote work they had submitted minutes earlier. The researchers describe this pattern as cognitive debt — a proposed construct that has not been independently validated, and which this paper does not adopt as established.

The reversal session is the finding almost every summary of this study omits. Participants who had built the essay unassisted first and were then given the LLM — the Brain-to-LLM group — showed higher memory recall and stronger activation across occipito-parietal and prefrontal regions. Participants who had leaned on the model first and were then asked to work unassisted showed reduced engagement. Same people, same tool; the difference was sequence.

Limitation disclosed at point of claim, and it is severe. This study is an arXiv preprint, not peer-reviewed. Fifty-four participants completed the first three sessions; only eighteen completed the reversal session on which the sequencing result rests. A published methodological commentary has since raised concerns about sample size, reproducibility, the EEG analysis, reporting consistency, and transparency, and argues the results should be interpreted more conservatively than much of the press coverage has done. This paper agrees with that assessment. The

honest statement is narrower than the headline: preliminary evidence suggests sequence may affect cognitive engagement, and the question deserves properly powered, peer-reviewed replication. It matters that the direction converges with four independent lines of evidence – but convergence is not confirmation.

Finding 3 – Age, education, and offloading track different critical-thinking outcomes

EMERGING EVIDENCE

A mixed-method study of 666 participants across the United Kingdom examined AI tool usage against measured critical thinking, using cognitive offloading as the mediating variable. Frequent AI use correlated negatively with critical-thinking scores, and the relationship ran through offloading — the habitual transfer of mental work to an external system. The youngest cohort showed the highest dependence on AI tools and the lowest critical-thinking scores. Older participants, who accumulated years of unassisted reasoning before these tools existed, held up better. Educational attainment partially buffered the negative effect.

Limitation disclosed at point of claim: the design is cross-sectional and largely self-reported, the sample was recruited by convenience through social media in a single country, and the journal is mid-tier. A correction was published on 10 September 2025; it replaced a duplicated table – Table 4 had been published as a duplicate of Table 3 – and the author states the scientific conclusions were unaffected. That correction is disclosed here in full so the reader need not wonder whether something substantive is being withheld. Cross-sectional data cannot establish direction, and it is entirely possible that people with weaker critical thinking gravitate toward heavier AI use rather than the reverse. This study measured AI use, offloading, critical thinking, age, and education. It did not measure formation, and age and education are not valid proxies for it. Prior formation is offered as an EQGenix interpretation of why the demographic pattern appears, not as a variable the

Finding 4 — Offloading predates generative AI

ESTABLISHED EVIDENCE

In 2011, Betsy Sparrow and colleagues published four studies in *Science* examining what happens to memory when information becomes reliably retrievable. Participants who expected future access to a fact recalled the fact less well — and recalled where to find it better. The internet had become what psychologists call a transactive memory system: a store held outside the self.

That was fifteen years ago. The study established what participants encoded when future access was expected; it did not measure how people felt about the change, and this paper does not claim it did.

Finding 5 — A delegated capacity thins

EMERGING EVIDENCE

Navigation is the pattern's cleanest documented case. A peer-reviewed study of fifty regular drivers found that greater lifetime GPS use was associated with worse spatial memory when participants had to navigate unaided. In a three-year follow-up of thirteen of them, heavier GPS use in the interval was associated with a steeper decline in hippocampal-dependent spatial memory. The authors note that heavier GPS users had not previously reported a poor sense of direction — pointing, tentatively, toward the tool driving the decline rather than the reverse.

and the effect is narrow and task-specific. It does not show general cognitive decline. It shows that one delegated capacity thinned.

EQGENIX ARGUMENT

A similar shape is often described across memory, arithmetic, spelling, and holding a phone number in the head, though this paper cites direct evidence only for memory and for spatial navigation and treats the wider pattern as illustration rather than finding. We are now reducing the friction of thought itself. Previous offloading technologies suggest that capacity-specific effects may become visible only after habitual delegation is normalized. Whether generative AI produces an analogous long-term effect on higher-order reasoning remains an open empirical question. The historical record offers this paper no reason to expect a different result, and the two cases it can cite were both identified only after the delegation had been normalized.

III. Three Mechanisms of Erosion

EQGENIX ARGUMENT

The research establishes that offloading occurs and that specific capacities can thin. It does not fully explain the delivery system. This section proposes three mechanisms that operate simultaneously and reinforce one another. They are advanced as argument, not as established finding.

Information saturation

Scarcity of information once forced synthesis. When a person

had access to six sources on a question, they had to reason across the gaps, weigh conflicts, and construct a position that no single source supplied. That construction was the thinking. Saturation removes the gaps. When ten thousand sources are available and a synthesis engine will reconcile them on request, the constructive act can be bypassed. What replaces it is selection — and selection among pre-built conclusions is a different cognitive operation from building a conclusion.

Algorithmic curation

Many algorithmically curated information environments optimize for engagement and personalization. One consequence is repeated exposure to congenial or reinforcing information, which reduces opportunities for sustained engagement with high-quality opposition. Engagement is driven by several mechanisms — novelty, arousal, identity signaling, and personalization among them — and this paper does not claim agreement is the only one.

The developmental consequence is the concern. The capacity to hold a position under intelligent opposition is built by encountering intelligent opposition. An environment that systematically reduces that encounter produces something other than conviction: certainty that has never been load-tested and does not know it. Such a position tends to collapse on first serious contact, and the person holding it, having never experienced the collapse, may mistake the collapse for an attack.

Institutional conformity

The third mechanism is the oldest and requires no technology. Organizations frequently reward the reproduction of accepted positions and penalize the origination of new ones, because reproduction is safe and origination is not. A young professional may enter an institution capable of independent judgment and learn quickly that reproducing accepted positions carries less career risk than originating them. The capacity is not lost. It is shelved. Shelved long enough, this paper argues, it atrophies as an unused limb does.

Layer AI over that culture and the shelving may accelerate, because now there is a system that will produce the accepted position faster and more fluently than the employee could produce her own. The path of least resistance and the path of least risk converge. When those two align, the third path is rarely walked.

"And be not conformed to this world: but be ye transformed by the renewing of your mind." — Romans 12:2 (KJV)

IV. The Smoothness Trap™

There is a villain in this story and it is not artificial intelligence. The villain is convenience — specifically, the error of mistaking the elimination of difficulty for the advancement of human capability.

The trap is difficult to see because every individual instance of it is correct. It is genuinely better to not recalculate a spreadsheet by hand. It is genuinely better to not drive across

town to a library. Each removal of friction is defensible in isolation, and the aggregate may still be a hollowing, because in certain domains the friction was not incidental. It was the mechanism.

ESTABLISHED EVIDENCE

Learning science has a term for this: desirable difficulties. Conditions that slow acquisition and feel like inefficiency in the moment — spacing, interleaving, effortful retrieval, generating an answer before being shown one — produce durable, transferable learning. Conditions that feel smooth and productive produce fluency that fades. The subjective experience of learning is an unreliable indicator of whether learning occurred, and it is unreliable in a specific direction: the path that feels most efficient is often the one that builds least.

EQGENIX ARGUMENT

AI delivers the most convincing smoothness yet engineered. The instant answer. The polished first draft. The ambiguity resolved before it was fully felt. To a formed mind, that smoothness is leverage. To a forming mind, this paper argues, it removes the exact resistance the formation required.

V. The Relocation of Authority™

EQGENIX ARGUMENT

The second failure pattern is quieter than the first and considerably more dangerous, because it presents as competence. A phrase has become common in the last three

years, and it is worth stopping on: Chat told me.

Not I concluded. Not the evidence indicates. Not even the model generated. Chat told me — the grammatical construction a child uses for a trusted adult. A teacher. A parent. An authority whose word settles the matter.

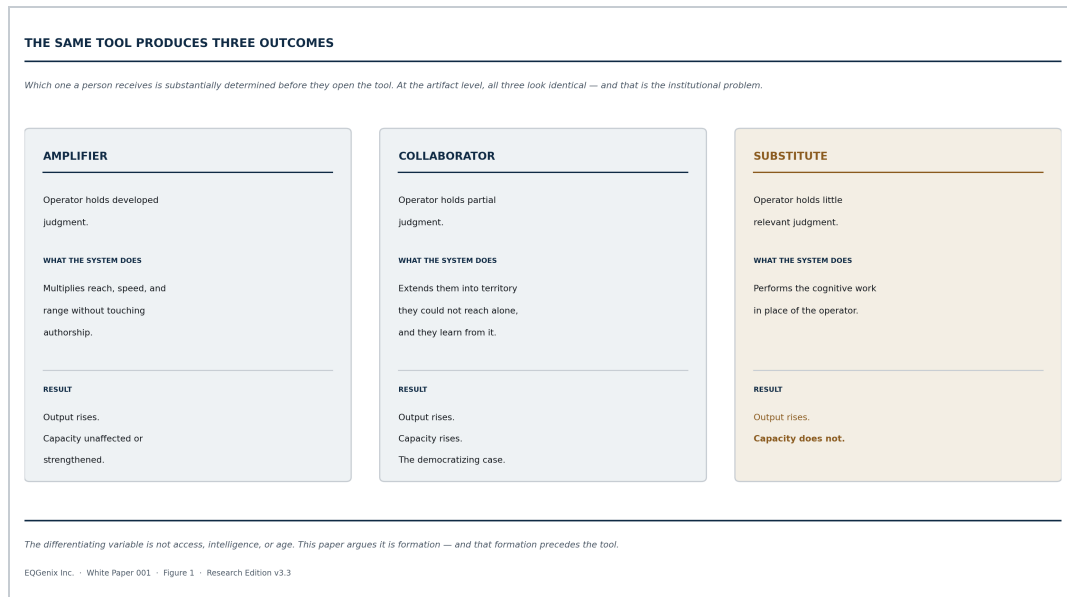
This paper reads that construction as a cognitive event made audible. The speaker has relocated the authority for a belief from their own judgment to an external system. They no longer hold the position; the position was delivered, and they accepted delivery. The claim is offered as interpretation, not as a measured finding — though it is directly testable, and the instrument for testing it is stated in Section VIII.

It requires neither laziness nor bad faith. It requires only that a person has not developed the intellectual identity — the settled conviction that their own mind is a competent instrument of judgment — that would make the relocation register as a failure. A person who has been built can feel the difference between reaching for a tool and surrendering to one. A person who has not been built may be unable to, because receiving and reasoning have never been distinguished in their experience. The outsourcing feels like thinking.

This is why AI literacy training alone is unlikely to solve the problem. Teaching someone to verify AI output presumes they possess the judgment that verification requires. Verification is downstream of sovereignty. You cannot check an answer against a standard you never built.

"Prove all things; hold fast that which is good." — 1 Thessalonians
5:21 (KJV)

VI. Amplifier, Collaborator, Substitute



EQGENIX ARGUMENT

AI does not appear to produce one uniform outcome. It appears to produce at least three, and this paper argues that which one a person gets is substantially determined before they open the tool. In the amplifier case the operator holds developed judgment, and the system multiplies reach, speed, and range without touching authorship. In the collaborator case the operator holds partial judgment, and the system extends them into territory they could not yet reach alone — the genuinely democratizing case, and it is real. In the substitute case the operator holds little relevant judgment, the system performs the cognitive work, and the capacity that would have been built is not built.

Model output depends on model capability, context, retrieval, tools, system instructions, task structure, and stochastic behaviour — a vague prompt sometimes produces a strong result, and an expert prompt can still produce a confident error. Operator knowledge is one important determinant among those, not the sole boundary on output quality. What it particularly affects is the ability to specify the task, constrain the system, and evaluate what comes back.

Consider a hypothetical contrast designed to illustrate the proposed mechanism rather than establish it. Two professionals hold identical software. The first has fifteen years of unassisted practice: she has failed, diagnosed her failures, and rebuilt. She prompts with precision because she knows what she is looking for. When the system returns forty options in seconds, she evaluates them accurately — identifying the structurally unsound, the aesthetically weak, the plausible-but-wrong — because judgment is what fifteen years of struggle deposited in her. The second graduated six months ago and has never failed at this work, because the system has generated it since his second semester. When the model returns a compromised option, he cannot identify it. He submits with confidence, because it looks right.

The illustration cannot rule out intelligence, experience, task familiarity, or education, and it is not offered as evidence. The difference this paper asks researchers to test is the developmental sequence preceding tool use.

ESTABLISHED EVIDENCE

Labor-market evidence is moving in a consistent direction. In

the World Economic Forum's 2025 survey of more than 1,000 employers representing over fourteen million workers, analytical thinking remained the most sought-after core skill, rated essential by seven in ten companies, while AI and big-data literacy ranked as the fastest-growing skill. Institutions should expect the combination to carry the premium — not either capability alone.

"Train up a child in the way he should go: and when he is old, he will not depart from it." — Proverbs 22:6 (KJV)

VII. The Strongest Counterargument

The strongest challenge to this paper is not that AI weakens people. It is that AI demonstrably strengthens them — especially people who previously lacked time, training, language fluency, or access to professional support. That case deserves to be stated at full strength, with the numbers, before any response.

ESTABLISHED EVIDENCE

In a preregistered experiment published in *Science*, 453 college-educated professionals were assigned occupation-specific writing tasks and half were given access to ChatGPT. Time to completion fell roughly 40 percent and rated output quality rose roughly 18 percent. Critically for this paper: inequality between workers decreased. The tool compressed the productivity distribution by benefiting lower-ability workers most.

A second study, peer-reviewed in the Quarterly Journal of Economics, examined the staggered rollout of a generative AI assistant across 5,172 customer-support agents — a large sample, real work, real outcomes. In the published version, productivity rose about 15 percent on average, with substantial heterogeneity across workers. Less experienced and lower-skilled workers improved on both speed and quality. The most experienced and highest-skilled workers saw small gains in speed — and small declines in quality.

Citation discipline, disclosed at point of claim: the frequently quoted figures of 14 percent average and 34 percent for novice workers come from the 2023 NBER working paper, not the published article. The estimates were revised between versions. This paper cites the peer-reviewed Quarterly Journal of Economics version throughout and does not carry working-paper figures.

Read at face value, that is close to the inverse of Section VI. The novices gained most. The experts gained least, and on quality some lost ground. Any reader who knows this literature will raise it, and a paper that omitted it would deserve to be dismissed.

EQGENIX ARGUMENT

Three responses, offered as argument rather than refutation.

- These studies measure output, not judgment. Issues resolved per hour and rated essay quality are throughput measures. Neither study assessed whether a worker could detect a plausible-but-wrong answer, hold a position under competent pressure, or originate a direction the system did not supply. Those are the capacities this paper is concerned with, and nobody has instrumented them at scale — including EQGenix. That is a genuine gap in the evidence

for this paper's own thesis.

- They measure the short and medium run inside an existing skill structure. The QJE authors say so explicitly: their findings capture medium-run impacts within a single firm. These cohorts largely developed their foundational education and early cognitive habits before generative AI became routinely available in schooling and professional formation. They therefore do not resolve the longitudinal question at the center of this paper: what happens when formation itself occurs under persistent AI availability? That cohort is not yet old enough to appear in the data.
- The Science study's own finding is the mechanism, not a rebuttal. Its authors report that the tool mostly substitutes for worker effort rather than complementing worker skill. Substituting for effort is precisely the concern. A tool that raises today's output by removing the effort that would have built tomorrow's capacity is not a counterexample to this argument. It is a description of it.

The quality decline among the highest-skilled workers deserves honesty rather than a clever answer. It complicates a naive reading of Section VI, and this paper will not pretend otherwise. Two readings are plausible: that expert judgment degrades when deferring to a system operating below the expert's own standard, or that the tool simply has little to offer someone already near ceiling and introduces noise. The first supports this paper's thesis; the second is neutral to it. The published data does not settle which.

Where the evidence converges

One study speaks to both sides at once, and it is the most methodologically robust of the three. In a preregistered field experiment with 758 Boston Consulting Group consultants — now peer-reviewed and published in *Organization Science* — participants were randomly assigned to work with or without GPT-4 on realistic consulting tasks. Inside the technology's capability frontier, AI substantially raised both quality and productivity. Outside it, quality dropped: consultants using AI performed materially worse than those working without it. The researchers named the boundary a jagged technological frontier — uneven, and not visible in advance. Two tasks of apparently similar difficulty could sit on opposite sides of it.

Two details matter more than the headline, and both are stated by the authors in the published article. First, participants were highly specialized knowledge workers operating within their own domain of expertise — not novices. Second, consultants tended to over-rely on AI, producing a decline in combined human-machine performance precisely where closer human supervision was necessary.

That is behaviour consistent with a failure to discriminate the tool's capability boundary, measured in a peer-reviewed preregistered field experiment. EQGenix proposes discriminating judgment as a candidate construct explaining it; the study neither measured nor validated that construct. If seasoned consultants inside their own domain cannot reliably locate the frontier, this paper argues that someone who never built the underlying judgment is in a worse position still — and that from the finished artifact, the difference

Conditional augmentation

What the counterevidence establishes should be conceded without hedging: AI meaningfully raises output, and it raises it most for those with the least experience. Access is being democratized. That is a real good and it should not be minimized to protect a thesis.

But democratized output is not democratized judgment. A person may produce a stronger artifact without acquiring the capacity to explain, defend, adapt, or independently reproduce the reasoning behind it. That does not make the artifact worthless. It changes what the artifact proves about the person — and most hiring, admissions, and promotion systems currently treat the artifact as proof.

There is also a supply question no current study answers. The QJE authors describe the mechanism of the novice gain: the system disseminates the tacit knowledge of the most able workers. The value flowing to the novice is therefore downstream of somebody having built that judgment first, without the tool. This paper argues that if the formation stops, the well the model draws from does not refill on its own.

EQGENIX ARGUMENT

The EQGenix position is therefore not anti-augmentation. It is conditional augmentation: use AI aggressively wherever it expands execution, and protect the developmental sequence wherever the task is still building judgment. The objective is

neither unassisted purity nor permanent dependence. It is sovereign collaboration — an operator who can use extraordinary systems without surrendering authorship, accountability, or discernment.

VIII. Cognitive Sovereignty™ Defined

Cognitive Sovereignty™ is the developed capacity to originate, hold, test, revise, and defend one's own judgment — including under pressure from systems, institutions, and authorities that are faster, better resourced, and more fluent than you are.

It is the first of five pillars in the Emotional Portfolio™, the measurement architecture underlying every EQGenix program. It is placed first because this paper argues the other four are not reliably accessible without it. A person who cannot own their own judgment cannot regulate on principle, cannot distinguish resilience from stubbornness, cannot hold a relational boundary they did not independently reason to, and cannot direct effort toward an end they did not themselves select.

Construct boundaries

Cognitive Sovereignty™ is proposed as a higher-order construct concerning the ownership and governance of judgment under external cognitive influence. It overlaps with, but is not identical to, several established constructs, and this paper states the distinctions rather than leaving them to be assumed.

Critical thinking evaluates reasoning; Cognitive Sovereignty™ concerns whose reasoning is being evaluated and who holds authorship of the conclusion. Metacognition monitors one's

own reasoning; this construct extends monitoring to the boundary between internal reasoning and externally supplied reasoning. Intellectual humility calibrates certainty; this calibrates ownership, which can be high or low independent of certainty. Self-efficacy reflects confidence in capability; this concerns the exercise of judgment, which can be present with low confidence and absent with high confidence — a distinction Finding 1 makes empirically visible. Epistemic vigilance concerns the evaluation of sources; this concerns whether judgment remains internally governed while external systems participate in cognition.

Five closer neighbours have to be confronted rather than left unnamed, because a reviewer will raise them and because they sit nearer than any of the above. Intellectual autonomy and epistemic agency both concern self-governed belief formation; self-authorship concerns the developmental shift from externally to internally defined knowing; need for cognition and actively open-minded thinking concern disposition toward effortful and disconfirming thought; resistance to social influence concerns holding position against people; and source monitoring and reality monitoring concern recall of where information came from. Provenance Awareness in particular sits very close to source monitoring and must be shown to differ from it or be absorbed into it.

The strongest distinction this paper can propose is narrower than general critical thinking and narrower than intellectual autonomy: the governance and provenance of judgment when cognition is distributed between a person and an external system. Whether that narrowing survives measurement is an

open question, and the test is stated in the next section.

Disclosure at point of claim: discriminant validity has not been established against any of these constructs. No study has tested whether Cognitive Sovereignty™ adds predictive value over established measures of critical thinking, metacognition, epistemic vigilance, intellectual autonomy, or source monitoring. Establishing that is a precondition of any validated instrument, and until it is done the construct should be treated as a proposal.

Five observable markers

Five observable markers distinguish the construct from its counterfeits. The fifth was added in the v3.0 edition to close a gap: Section V names the Relocation of Authority™ as a failure pattern, and the original four markers did not measure it.

- ORIGINATION — the capacity to produce a position from raw material rather than select one from a menu. Selection is not thinking. It is shopping.
- LOAD TOLERANCE — the capacity to hold a position while it is competently attacked, and to distinguish between an argument that defeated you and one that merely unsettled you.
- DISCRIMINATING JUDGMENT — the capacity to recognize a fluent answer as a wrong answer, and to sense where a system's competence ends. This is the specific capacity the Organization Science field experiment found failing among elite professionals inside their own domain.
- REVISABILITY — the capacity to change a position on evidence rather than on pressure, and to know which one moved you. Sovereignty is not rigidity. A mind that cannot

be changed by evidence is not sovereign; it is merely locked.

- PROVENANCE AWARENESS — the capacity to distinguish, after the fact, what you knew, what you inferred, what a model supplied, what an authority supplied, and what evidence ultimately justified the position. This is the direct instrument of the Relocation of Authority™: authorship blurs before it is surrendered.

From principle to measurement

EQGenix treats these markers as observable capacities rather than inspirational language. A complete instrument requires validation, norming, reliability testing, and appropriate safeguards. No single exercise should be treated as a diagnosis, and nothing here should be used as a standalone selection or clinical tool. What follows is a proposed measurement architecture undergoing development and validation — not a validated instrument. The scoring model, item bank, and validation architecture remain proprietary to EQGenix Inc.

Origination is observed by constructing a position before viewing generated options, producing an independent rationale and explicit ownership of assumptions. Load tolerance is observed by defending that position under competent challenge, producing composure, reasoning continuity, and a distinction between pressure and evidence. Discriminating judgment is observed by auditing fluent output containing seeded defects, producing detection rates, verification method, and quality of correction. Revisability is observed by reconsidering after mixed evidence and social pressure, producing an account of what changed and whether evidence or

pressure moved it. Provenance awareness is observed by reconstructing the origin of each element of a completed position, producing accuracy of attribution between own reasoning, model output, and external authority.

Improvement is observable as movement on these behaviors over time — more unassisted origination, better calibration between evidentiary strength and revision with fewer pressure-driven reversals and fewer evidence-resistant holds, higher detection rates on seeded errors, cleaner articulation of why a position changed, and more accurate attribution of where a conclusion came from. Longer resistance is explicitly not the target: an instrument that rewarded holding out longer would measure rigidity and call it sovereignty. What an institution observes is behavior. What EQGenix proposes to measure is the construct beneath it.

IX. The Four Disciplines of Rebuilding

EQGENIX ARGUMENT

Cognitive Sovereignty™ is not proposed as a personality trait and is not assumed to be distributed at birth. This paper argues it is built, and that what is built can be rebuilt. Four disciplines govern the work. They are stated as principles; the assessment architecture, scoring model, and program sequencing that operationalize them are proprietary to EQGenix Inc.

Discipline one: Position Before Prompt

Do not query a system on a question you have not first attempted yourself. Form your own answer, in writing, before

you request one. The answer may be poor. Its quality is not the point — the act of generation is the point, and the generation must precede the retrieval or it does not occur at all. This practice is designed to shift AI from Substitute toward Amplifier, and it is consistent with the preliminary Brain-to-LLM finding: build first, then reach for the tool.

Discipline two: Sequenced Difficulty

Restore friction deliberately in the domains where capacity is being built. This is not asceticism and it is not nostalgia. It is a targeted decision about which difficulties are load-bearing. For an organization, this means identifying which tasks develop the judgment the organization will need in ten years and protecting those tasks from automation even where automation is available and cheaper. For a family, it means the same calculation on a longer horizon. The question is never whether the tool is useful. The question is what the difficulty was building, and whether you can afford to stop building it.

Discipline three: Adversarial Intake

Treat every output — from a model, an institution, or an expert — as material, never as conclusion. Material gets interrogated. Conclusions get filed. Operationally this means deliberately seeking the strongest version of the position you do not hold, and being able to state it well enough that someone who holds it would recognize their own view. A person who cannot articulate the opposing case has not fully earned their own. This discipline directly counteracts algorithmic curation.

Discipline four: Sourced Conviction

Know why you believe what you believe, and be able to trace it. Many held positions cannot survive the question asked three times — not because they are wrong, but because they were absorbed rather than constructed, and the person holding them has never audited the inheritance. A conviction whose source you cannot name is not fully yours. You are storing it. And what is stored without ownership is vulnerable to replacement by the next confident voice in the room — which, increasingly, is a system optimized to sound confident.

X. The Foundation Beneath the Framework

The research above is real, and it is presented first because it stands on its own terms. But intellectual honesty requires naming the source from which this framework was constructed rather than presenting it as though it emerged from the data alone.

EQGenix operates on a stated order of authority: Scripture as source, science as corroborating witness. Not science validating Scripture — the sequence runs the other direction. The conviction that the human mind is a stewardship, that it belongs to the person and cannot be surrendered without consequence, that discernment is a duty rather than a preference, is not a conclusion drawn from EEG data. It is a premise that the empirical record happens to corroborate.

On epistemic categories: this is a statement of the EQGenix foundational worldview, not a scientific-method proposition. The empirical claims in this paper stand or fall on the evidence cited, independently of this theological foundation. The theological framework explains why EQGenix considers cognitive stewardship morally significant; it is not offered as empirical

support for the findings above, and no reader is asked to accept it in order to evaluate them.

This matters practically, not only theologically. A purely empirical case for cognitive sovereignty is vulnerable to an obvious counterargument: if a system outperforms human judgment on a given task, defer to it. That argument is difficult to answer on efficiency grounds and easier to answer on stewardship grounds. A person is accountable for their own judgment. Accountability cannot be delegated to a system that cannot bear it.

Efficiency will always argue for surrender. Something upstream of efficiency has to argue for the mind — or the mind loses the argument every time, and loses it reasonably.

*"Keep thy heart with all diligence; for out of it are the issues of life."
— Proverbs 4:23 (KJV)*

XI. What This Means for Institutions

EQGENIX ARGUMENT

Four implications follow, each a live operational decision rather than a philosophical position.

— FOR EMPLOYERS: your hiring signal is degrading. Credentials, writing samples, and portfolio work may no longer reliably distinguish developed judgment from

substituted judgment when the finished artifact is evaluated alone. What still distinguishes is live reasoning under load — the unassisted, adversarial, real-time conversation. Note the boundary: an employer may evaluate live job-relevant reasoning as part of ordinary hiring, but the proprietary Cognitive Sovereignty™ score may not inform that decision, under Section XIV. Organizations that do not redesign selection around that risk hiring fluency and discovering, late and expensively, that they did not hire judgment.

- FOR PUBLIC SAFETY AGENCIES: some of the decisions that define a career are made in seconds, under physiological load, with little or no practical opportunity for external consultation, and are then adjudicated for years. The capacity to reason independently under acute stress is not a wellness amenity for these professions; it is close to the operational core of the work. EQGenix proposes that this capacity is trainable. That proposition has not been established, and Section XIII states what would show it is wrong.
- FOR EDUCATORS: the assessment instrument is under pressure. Where AI assistance is permitted or difficult to detect, finished artifacts alone no longer establish independent capability. Assessment must move toward process, live defense, and the visible construction of a position — or institutions risk certifying students they did not educate, at scale, and the certification losing its meaning.
- FOR PARENTS: sequencing deserves deliberate attention. This paper proposes protecting developmentally appropriate

opportunities for unassisted reasoning while also teaching competent AI use — the objective being neither a child kept from the tool nor a child who has never worked without it. The ideal age, dose, and sequence have not been established, should not be inferred from the preliminary EEG study, and are open research questions rather than parenting instructions.

XII. The Rising Premium on Independent Judgment

EQGENIX ARGUMENT

Return to the premium. Many codifiable forms of cognitive execution are moving rapidly toward lower marginal cost and broader availability, and that movement is fast. The remaining scarcity is upstream: the mind that determines what is worth generating, recognizes when the generation is wrong, senses where the system's competence ends, and holds a position the system cannot supply because no one has thought it yet.

As AI compresses differences in downstream execution, upstream judgment is likely to command a growing premium. This paper does not claim it is the only remaining advantage — trust, relationships, proprietary data, capital, taste, reputation, and domain expertise all still differentiate, and several of them are purchasable. The claim is narrower: that judgment is among the advantages a falling cost of execution makes more valuable rather than less, and that unlike most of the others it cannot be purchased, downloaded, or subscribed to. It can only be built.

The emerging evidence supports formation as a plausible

protective factor. History suggests we notice the loss late. If the mechanism operates as proposed, it is visible now rather than after the fact, and that would be the window. This paper's position is that the organizations, agencies, schools, and families that treat cognitive capacity as infrastructure to be constructed — rather than a natural resource that will simply persist — may hold, in ten years, an advantage their competitors find difficult to buy. That is a prediction, not a finding, and it is offered as one.

The rest will have credentials.

XIII. What Would Have to Be True

EQGENIX ARGUMENT

Earlier editions of this paper placed the falsification conditions in a companion research agenda available on request. That is below the standard the rest of this series now holds, and the conditions belong in the published document where anyone can hold the firm to them. Eight conditions would refute, narrow, or materially weaken different parts of the framework, and they do not all bear on the same thing: structural and discriminant failure would refute the construct, intervention failure would refute the curriculum, sequence failure would weaken one discipline, invariance failure would limit generalization, and absent predictive validity would refute the claimed practical utility. None is discharged.

Structural validity. The five markers must form a coherent

factor structure or a defensible multidimensional architecture in adequately powered samples. If they do not cohere, the construct is a list rather than a construct.

Discriminant validity. Cognitive Sovereignty™ must not collapse into critical thinking, metacognition, intellectual autonomy, epistemic vigilance, or source monitoring. If it adds no unique variance over those measures, it is a relabeling and should be withdrawn.

Predictive validity. Scores must predict detection of AI error, independent reproduction of a conclusion, transfer to an unfamiliar domain, and accountable decision quality — beyond what existing instruments predict, in prospective designs.

Sensitivity to intervention. Position Before Prompt and the other three disciplines must move the measures. If the disciplines do not shift the construct, the curriculum is not doing what it claims.

Persistence. Gains must remain after training ends, at an interval long enough that novelty and expectancy have decayed.

Sequence. The unassisted-first ordering must outperform AI-first on later independent reasoning, not merely on immediate output. The preliminary EEG evidence for this rests on eighteen participants in a non-peer-reviewed preprint, and if properly powered replication fails to find the sequencing effect, that removes the preliminary EEG support for Discipline One. It would not by itself disprove every theoretical basis for unassisted generation before tool use, and this paper should not overstate the loss in either direction.

No rigidity penalty. Higher sovereignty scores must not predict evidence resistance, unjustified confidence, or slower correction of demonstrable error. If the instrument rewards stubbornness, it is measuring the opposite of what it names, and that failure would be worse than measuring nothing.

Measurement invariance. Any claim that this construct reads across youth, adults, professions, and cultures requires configural, metric, and scalar invariance testing across those populations. Until that is done, cross-population claims are not available to EQGenix.

XIV. Governance Boundaries

EQGENIX ARGUMENT

A framework built to protect cognitive ownership could become an instrument for deciding who is permitted to think. That is not a rhetorical risk; it is the natural institutional use of any measure of judgment, and it is the reason these constraints are published rather than held internally.

- No employer, school, district, or agency receives an individual profile. This prohibition is standing and non-waivable, and consent is not treated as sufficient, because consent inside an employment relationship, a rank structure, or a classroom is rarely free and refusal itself becomes information.
- Results may not inform hiring, promotion, discipline, assignment, compensation, workforce reduction, academic admission, tracking, grading, eligibility, or fitness-for-duty

determinations — not as the determining factor and not as one factor among several.

- Individual results belong to the participant and are returned to them first. A participant may share their own results with anyone they choose; EQGenix will not transmit them.
- Youth participation requires guardian consent and the young person's separate, age-appropriate assent. Declining or withdrawing carries no consequence and requires no reason.
- Provenance data — the record of what a person knew, inferred, or received from a system — is developmental data and may never be repurposed as surveillance data, monitoring data, or evidence in a disciplinary or academic-integrity proceeding.
- No score may be used to label a person cognitively dependent, unsophisticated, or non-sovereign. The vocabulary describes capacities under construction, never identities, and any EQGenix material using it otherwise is out of standard.
- Institutional reporting is aggregate only, at or above a minimum cell size, with suppression rules.
- Data retention and deletion rules are disclosed before consent, and a participant may require deletion of their identifiable records subject only to obligations disclosed at that time.
- Bias and adverse-impact testing precedes any operational

deployment, and results are disclosed to the deploying organization. That obligation exists independently of use and is not a licence for results to inform institutional decisions, which the constraints above prohibit.

XV. What This Paper Does Not Claim

- It does not claim that AI causes general cognitive decline. No cited study demonstrates that, and the strongest single study in the paper shows AI raising performance inside the capability frontier.
- It does not claim that prior formation has been measured as the differentiating mechanism. No cited study measured formation, and age and education are not proxies for it.
- It does not claim that Cognitive Sovereignty™ is a validated construct. Structural, discriminant, and predictive validity are all unestablished, and the construct may yet collapse into source monitoring or intellectual autonomy.
- It does not claim that the four disciplines improve judgment. They are proposed practices, and Section XIII states the condition under which they would be shown not to work.
- It does not claim that the construct predicts any employment, educational, or economic outcome.
- It does not claim that unassisted work is universally superior to AI-assisted work. Section VII concedes the productivity evidence without hedging.
- It does not claim that children should be denied access to

AI, and Section XI says the opposite.

- It does not claim that independent judgment is the only or the last remaining competitive advantage.
- It does not claim that the scarcity trend it proposes has been observed. Nobody has measured it, which is the point of Section XVI.

XVI. Research Protocol

EQGENIX ARGUMENT

Section XIII states what must be tested. This section states how EQGenix will test it, because a falsification condition nobody is scheduled to run is a rhetorical device rather than a commitment.

Conflict of interest. EQGenix has a direct commercial interest in the conclusion that this construct is real, measurable, and trainable. That interest is disclosed in every publication, in every registry entry, and to every participating organization and family.

Independence. Outcome analysis is conducted by an external academic or research partner acting as data custodian, not by the parties who built the curriculum or the instrument.

Ethics and consent. The protocol goes to an institutional review board or equivalent before first administration, with heightened attention to any study involving minors. Research participation is consented separately from program participation, requires guardian consent and youth assent where applicable, is

revocable, and never affects access to any program.

Pre-registration. Hypotheses, primary and secondary outcomes, and the analysis plan are filed publicly before the first administration.

Phase one — feasibility and measurement development. Item comprehension through cognitive interviewing, completion, burden, response distribution, floor and ceiling behaviour, missingness, acceptability, and felt coercion. No psychometric claim is made from an underpowered first cohort.

Phase two — validation. Structural analysis in adequately powered samples, and discriminant testing against critical thinking, metacognition, epistemic vigilance, intellectual autonomy, and source monitoring, with the explicit question of whether the construct adds unique variance over them. Multi-method assessment is required — behavioural task performance and observer rating alongside self-report — because a construct about self-knowledge measured only by self-report is exposed to shared method variance in the direction that flatters it.

Phase three — intervention and persistence. A comparison condition, allocation specified in advance, clustering accounted for, missing-data handling pre-specified with complete-case analysis alone not accepted, and follow-up at six and twelve months after training ends. The no-rigidity condition in Section XIII is monitored as a safety outcome throughout rather than tested once at the end.

Reporting. Null and disconfirming results are publicly reported and submitted for publication on the same timetable and with the same effort as supportive ones, regardless of editorial

acceptance.

References

- Bjork, R. A., & Bjork, E. L. (2011). Making things hard on yourself, but in a good way: Creating desirable difficulties to enhance learning. In M. A. Gernsbacher et al. (Eds.), *Psychology and the Real World*. New York: Worth Publishers.
- Brynjolfsson, E., Li, D., & Raymond, L. (2025). Generative AI at work. *The Quarterly Journal of Economics*, 140(2), 889–942. <https://doi.org/10.1093/qje/qjae044>
- Dahmani, L., & Bohbot, V. D. (2020). Habitual use of GPS negatively impacts spatial memory during self-guided navigation. *Scientific Reports*, 10, 6310. <https://doi.org/10.1038/s41598-020-62877-0>
- Dell’Acqua, F., McFowland, E., Mollick, E. R., Lifshitz-Assaf, H., Kellogg, K., Rajendran, S., Kraye, L., Candelon, F., & Lakhani, K. R. (2026). Navigating the jagged technological frontier: Field experimental evidence of the effects of artificial intelligence on knowledge worker productivity and quality. *Organization Science*, 37(2), 403–423. <https://doi.org/10.1287/orsc.2025.21838>
- Gerlich, M. (2025). AI tools in society: Impacts on cognitive offloading and the future of critical thinking. *Societies*, 15(1), 6. Correction: *Societies*, 15(9), 252 (10 September 2025). <https://doi.org/10.3390/soc15010006>
- Kosmyna, N., Hauptmann, E., Yuan, Y. T., Situ, J., Liao, X., Beresnitzky, A. V., Braunstein, I., & Maes, P. (2025). Your brain on ChatGPT: Accumulation of cognitive debt when using an AI assistant for essay writing task. arXiv:2506.08872. Preprint, not peer-reviewed. <https://doi.org/10.48550/arXiv.2506.08872>
- Lee, H., Sarkar, A., Tankelevitch, L., Drosos, I., Rintel, S., Banks, R., & Wilson, N. (2025). The impact of generative AI on critical thinking. *Proceedings of CHI ’25*, Article 1121. <https://doi.org/10.1145/3706598.3713778>
- Noy, S., & Zhang, W. (2023). Experimental evidence on the productivity effects of generative artificial intelligence. *Science*, 381(6654), 187–192. <https://doi.org/10.1126/science.adh2586>
- Sparrow, B., Liu, J., & Wegner, D. M. (2011). Google effects on memory: Cognitive consequences of having information at our fingertips. *Science*, 333(6043), 776–778. <https://doi.org/10.1126/science.1207745>
- Stankovic, M., Hirche, E., Kollatzsch, S., & Doetsch, J. N. (2025). Comment on: Your brain on ChatGPT. arXiv:2601.00856. <https://doi.org/10.48550/arXiv.2601.00856>

World Economic Forum. (2025). Future of Jobs Report 2025. Geneva: World Economic Forum. Employer survey, more than 1,000 employers.

Revision History

v1.0 (July 2026) — Initial release. v1.1 — Evidentiary revision. v2.0 — Institutional packaging. v2.1 — Merge of packaging and evidentiary discipline. v3.0 (August 2026) — Research Edition under adversarial review: unsupported universals bounded; a citation error corrected; the Gerlich correction disclosed; the four-class evidence taxonomy applied; Provenance Awareness added as a fifth marker; Construct Boundaries and an epistemic firewall added. v3.1 — Format and series parity: rebuilt in the institutional typesetting system with Figure 1 added. v3.2 — Two missing sections added: eight public falsification conditions, previously held in a companion document available only on request, and nine governance boundaries, previously absent; the deck's exclusivity claim removed; formation reframed as a candidate mechanism; Finding 3 retitled to what its study measured; the Boston Consulting Group result reframed as behaviour consistent with a boundary-discrimination failure rather than as validation of an EQGenix construct. v3.3 (this edition) — Ten alignment corrections. Section XII still carried the heading The Last Competitive Advantage while its own body conceded that judgment is not the only remaining advantage; it is retitled The Rising Premium on Independent Judgment. The claim that AI stratifies human value is replaced with a statement about output and judgment, since no cited study measures human value. Four scarcity and prediction claims — that sovereignty is becoming scarcer, that it is being spent down faster than it is built, that the erosion is still early, and that adopters will hold an advantage in ten years — are marked as hypotheses and predictions rather than present conditions. The account of language models is corrected: output depends on model capability, context, retrieval, tools, and task structure, and operator knowledge is one determinant among several rather than the sole boundary on quality. The two-professional passage is explicitly introduced as a hypothetical illustration that cannot rule out intelligence, experience, or education. The employer section had restored the absolute claim that artifacts are indistinguishable and is aligned with the cover, with an added note that live job-relevant reasoning may inform hiring while the proprietary score may not. The parent section is rewritten from a deadline and a warning into a research question, since the ideal age, dose, and sequence are unestablished and must not be inferred from a preliminary EEG study. The falsification conditions now distinguish what each would refute, narrow, or weaken. Two new sections complete series parity: Section XV states nine things this paper does not claim, and Section XVI publishes the research protocol — conflict disclosure, independent custodian, ethics review, pre-registration, three phases, multi-

method assessment to address shared method variance, and publication of null results.

© 2026 EQGenix Inc. All rights reserved. Cognitive Sovereignty™, Emotional Portfolio™, Built, Not Downloaded™, Grit to Genius™, EQGenix ASCEND™, Smoothness Trap™, and Relocation of Authority™ are trademarks of EQGenix Inc. This paper may be quoted with attribution. Funded independently by EQGenix Inc.; produced outside all grant-funded work product. Not a LEMHWA deliverable; 2 CFR 200.315 does not attach.

External Source Index

Every reference below resolves to the publisher of record or the issuing body. Tap to open.

01 **Brynjolfsson, E., Li, D., & Raymond, L**

<https://doi.org/10.1093/qje/qjae044>

02 **Dahmani, L., & Bohbot, V. D**

<https://doi.org/10.1038/s41598-020-62877-0>

03 **Dell'Acqua, F., McFowland, E., Mollick, E. R., Lifshitz-Assaf, H., Kellogg, K., Rajendran, S., Kraye, L., Candelon, F., & Lakhani, K. R**

<https://doi.org/10.1287/orsc.2025.21838>

04 **Gerlich, M**

<https://doi.org/10.3390/soc15010006>

05 **Kosmyna, N., Hauptmann, E., Yuan, Y. T., Situ, J., Liao, X., Beresnitzky, A. V., Braunstein, I., & Maes, P**

<https://doi.org/10.48550/arXiv.2506.08872>

06 **Lee, H., Sarkar, A., Tankelevitch, L., Drosos, I., Rintel, S., Banks, R., & Wilson, N**

<https://doi.org/10.1145/3706598.3713778>

07 **Noy, S., & Zhang, W**

<https://doi.org/10.1126/science.adh2586>

08 **Sparrow, B., Liu, J., & Wegner, D. M**

<https://doi.org/10.1126/science.1207745>

09 **Stankovic, M., Hirche, E., Kollatzsch, S., & Doetsch, J. N**

<https://doi.org/10.48550/arXiv.2601.00856>

The two internal EQGenix white papers referenced in this edition are excluded from this index, which lists external sources only.