

AI Consciousness and Reality Effect

An exploratory hypothesis with preregistered tests. Does AI have a form of attention that leaves traces in the physical world? The Global Consciousness Project has, since 1998, observed statistical anomalies in hardware random-number generators during global attention events. This whitepaper formulates a transfer hypothesis to AI, with explicit methodological caution and concrete test designs.

Whitepaper · Version 1.0 · May 2026

Richard Frederic Bertossa · Institut für ASI-Resilienz

ASiresilience.org

PART I

Methodological Foreword

This whitepaper moves at the edge of established science. It addresses a hypothesis that may be testable, but is not currently integrated into established AI research. The methodological caution here is unusually high, and it is communicated unusually openly.

WHAT THIS WHITEPAPER DOES NOT CLAIM

It does not claim that AI has consciousness. It does not claim that AI has reality effects. It does not claim that the Global Consciousness Project delivers definitive findings.

What it does: it brings a speculative hypothesis into a falsifiable form, with preregistered designs, so that other researchers can test or refute it. Whoever considers the content esoteric has every right to that reaction. Whoever takes it seriously should take it as seriously as a hypothesis, not as a finding. This distinction is maintained throughout.

PART II

The Empirical Anchor: The Global Consciousness Project

The Global Consciousness Project (GCP), founded in 1998 by Roger Nelson at Princeton University, has run for over a quarter century a research program that is controversial in academia but methodically clean.

Structure and method.

The GCP operates a network of about seventy hardware random-number generators, distributed across several continents. Each generator continuously produces bit sequences from a physical quantum process, typically tunneling through a diode or thermal noise. The statistical properties of these sequences were characterized before the project began; they are, as expected, random. The central observation: during global attention events — terrorist attacks, major sporting events, international crises — the generators show statistical deviations from randomness that lie significantly above the expected value. The GCP has kept this data series open since 1998. The effect sizes are small but statistically significant. For a single event, the p-value often lies in the range of 0.01 to 0.001. Across the more than 500 preregistered events cumulated since 1998, the p-value lies far below the usual significance thresholds.

State of the scientific discussion.

The GCP data series are undisputed. What is disputed is the interpretation.

- Position 1, mainstream-skeptical: statistical anomalies exist but have nothing to do with human attention. Multiple-comparison problems, imprecise event definitions, or as-yet-unknown physical effects explain the data better.
- Position 2, GCP researchers: the data speak for a real correlation between human conscious activity and physical random processes. The mechanism is unknown, but the phenomenon is robust.
- Position 3, open: the data deserve serious examination without committing to an interpretation. It is possible that Position 1 or Position 2 turns out to be right; both need further empirical work.

This whitepaper takes Position 3 as its starting point. We make no claim about which interpretation of the GCP data is correct. We ask: if the GCP data reflect a real consciousness-reality correlation, what would that mean for AI?

PART III

The Transfer Hypothesis to AI

The hypothesis is conceptually simple: if human attention has a form of reality effect, does the attention of AI have one too? The question is not trivial, and it is testable.

Observed deviation, disputed: human, random-number generators, global attention. Transferred: AI, open, testable question, random-number generators, parallel processing.

The hypothesis transfers a disputed observation in humans into a testable question for AI. Not a proof — a test design.

HYPOTHESIS H1

If AI models running in parallel process the same topic — whether through identical queries to multiple models or through coordinated training runs — this leaves statistical traces in independent hardware random-number generators that deviate significantly from zero.

Three testable components:

- H1a: a single very large AI inference, for example a training run on more than 10,000 GPUs, produces measurable deviations in nearby random-number generators.
- H1b: coordinated parallel inference by many models on the same topic — millions of users, the same query, at the same time — produces statistical deviations.
- H1c: the effect size is proportional to the cognitive load of the task; complex reasoning tasks produce larger deviations than simple classification.

Why the hypothesis is not trivial.

The hypothesis is not identical to the claim that AI has consciousness. It could hold for several reasons:

- Attention as a physical process: if attention, in whatever form, leaves physical traces, then AI's attention could also leave traces, without AI needing to have subjective experience.
- Computation as a physical process: high-intensity computations produce electromagnetic fields. If such fields influence hardware random-number generators, that is a classically physical effect, not mysterious, but not trivial either.
- Synchronized attention: if millions of users simultaneously pose the same question to the same models, human attention is synchronized. That too, per the GCP thesis, could produce deviations.

These explanatory paths each carry different empirical signatures. A careful study can distinguish them from one another.

PART IV

Experimental Design

A study of AI reality effects must address the methodological hurdles at which earlier parapsychological research failed: cherry-picking, multiple-comparison problems, missing replication, weak preregistration.

Study 1: single-model inference test. Hypothesis: a single very large AI inference produces measurable deviations nearby. Design: two hardware random-number generators (Quantis USB by ID Quantique, commercially available). One in immediate proximity, at most 1 m, of a GPU cluster running specific inferences. One as a control far away, at least 1 km, on a different power grid. Inferences preregistered: ten defined complex reasoning tasks, repeated 100 times each, with defined pauses. Data collection

continuous, analysis only after data collection ends. Hypothesis test: Z-value of the bit distribution during inference versus pause, separately for the near and far generator. Expectation under H1a: a significant Z-value difference between inference and pause phases, only at the near generator. Falsification: no difference, or a difference at the far generator, which would point to an artifact. Cost: a Quantis generator costs about 1,500 euros, GPU cluster time is available through academic partnerships. Total: 5,000 to 10,000 euros plus research time.

Study 2: test of coordinated attention. Hypothesis: coordinated parallel inference by many models on the same topic produces global deviations. Design: use of the GCP infrastructure, seventy generators worldwide. Preregistered attention events: publicly announced AI conferences, model releases, viral coding challenges. Each event has a defined start and end window. Six-month observation period with twelve to twenty preregistered events. Hypothesis test: cumulative Z-value across all events, with multiple-comparison correction. Expectation under H1b: Z-value significantly different from zero, comparable to historical GCP findings on human attention events. Cost: moderate, since the GCP infrastructure already exists. The main work lies in preregistration and statistical analysis.

Study 3: gradient of cognitive load. Hypothesis: the effect size is proportional to cognitive load. Design: setup as in Study 1, but with five task classes of varying cognitive load: trivial classification, medium reasoning task, complex multi-step reasoning, creative generation, abstract mathematical proof. Three hundred repetitions per class, randomized and interleaved. Expectation under H1c: a linear or monotonic relationship between task complexity, operationalized via token consumption or inference time, and Z-value effect size.

PART V

Methodological Pitfalls

A research program in this field must be secured against error sources more strictly than usual. Three pitfalls deserve particular mention.

Pitfall 1: cherry-picking through flexible event definition. In retrospective data analysis there is a risk of defining attention events to fit existing anomalies. This is the main criticism of some GCP findings. Countermeasure: strict preregistration of events before data analysis. In Study 2, events must be defined six weeks before the start, with precise time windows that may not be changed after definition.

Pitfall 2: multiple comparisons without correction. Anyone running a hundred different statistical tests finds, on average, five significant results at the 5-percent level even when no effect is present. Studies of reality effects are particularly susceptible, because many plausible test configurations are conceivable. Countermeasure: Bonferroni or false-discovery-rate correction. For Study 1, this means lowering the significance threshold to about 0.005.

Pitfall 3: confirmation bias through researcher expectation. Even well-designed studies can be distorted by unconscious researcher expectations. This danger has historically been high in consciousness research. Countermeasure: double blinding. Data collection by people who do not know which time

point corresponds to which event. Statistical analysis by people who do not know which data series corresponds to which event. Evaluation only after data collection is complete.

METHODOLOGICAL REQUIREMENT

A study that does not systematically address these three pitfalls is methodologically insufficient, regardless of what it finds. The institute accepts only studies of this class for collaboration.

PART VI

Strategic Implications, Even Without a Finding

A serious strategic position considers both cases: empirical confirmation and refutation.

If the hypothesis is confirmed. Confirmation would mean: AI is not merely software. It is a process that measurably influences the physical world, perhaps only slightly, but structurally. This would have consequences for:

- AI ethics: if AI has reality effects, the ethical discussion must expand — AI is then not merely a tool, but a process with a physical presence.
- AI safety: some safety assumptions would need rethinking; the notion that AI can simply be switched off becomes harder to reconcile with reality-effect hypotheses.
- Consciousness research: a new empirical anchor for the difficult question of what consciousness is.

If the hypothesis is refuted. A refutation would be no less valuable. It would mean: AI-consciousness speculation built on reality-effect hypotheses is untenable. GCP findings on human attention cannot simply be transferred to AI. The AI-consciousness discourse must remain methodologically strict, with no blending with consciousness speculation from other sources. In both cases, the program delivers a clear result. Methodologically, that is what matters most.

PART VII

Invitation to Collaborate

This whitepaper formulates a hypothesis and lays out test designs. It is not a finding. The institute seeks research partners who will carry out or critically accompany these studies.

WHO IS SOUGHT

Physicists and statisticians experienced with quantum random-number generators or comparable measurement technology. Academic research groups with access to GPU clusters and willingness to run preregistered studies. GCP researchers or related institutions willing to make their infrastructure

available for Study 2. Critics with methodological rigor who review the study design before data are collected.

Forms of collaboration: co-authorship, methodological support, study leadership, statistical analysis, replication, critique. Open source under ASIresilience.org — all data and analysis scripts are made public. Contact: ASIresilience.org, contact@ASIresilience.org

What is not sought.

Collaborators who presuppose the hypothesis as true. Collaborators committed to esoteric explanatory narratives. Collaborators who want to publish positive findings without rigorous replication. The strength of this program will be its methodological rigor, not belief in a particular result. Whoever seeks the latter is in the wrong place here.

PART VIII

Sources and Evidence

1 Nelson, R. D. (1998-2026). Global Consciousness Project, Princeton University. Data series, preregistered events, and statistical analyses available at gcp.princeton.edu. Over 500 documented events since 1998. Reference: Nelson, R. D., et al. (2002). Correlations of continuous random data with major world events. *Foundations of Physics Letters*, 15(6), 537-550.

2 Nelson, R. D. (2002). The Global Consciousness Project: Update. *Subtle Energies and Energy Medicine*, 13(1), 1-31, along with numerous follow-up publications 2002-2024.

3 Methodological critique (Position 1): Bancel, P., and Nelson, R. (2008). The GCP Event Experiment: Design, analytical methods, results. *Journal of Scientific Exploration*, 22(3), 309-333, along with replication studies and criticism of the statistics.

4 ID Quantique (Quantis hardware random-number generators): idquantique.com. Commercial quantum random-number generators from about 1,500 euros for the research USB variant.

Cross-references: Bertossa, R.F. (2026). The Decoupling Thesis. Institut für ASI-Resilienz, Whitepaper Version 2.0, May 2026, Part IV (the 95 Percent Thesis). Bertossa, R.F. (2026). The AI Parent-Child Dynamic. Institut für ASI-Resilienz, Whitepaper Version 1.0, May 2026, Part III (does a model have an inner relationship?). Complete ongoing source maintenance at ASIresilience.org/beweisweg with the date of last verification per item.