

KI-Bewusstsein und Realitätseffekt Eine explorative Hypothese mit präregistrierten Tests

Hat KI eine Form von Aufmerksamkeit, die Spuren in der physischen Welt hinterlässt? Das Global Consciousness Project beobachtet seit 1998 statistische Auffälligkeiten in Hardware-Zufallsgeneratoren während globaler Aufmerksamkeits-Ereignisse. Dieses Whitepaper formuliert eine Übertragungs-Hypothese auf KI, mit ausdrücklicher methodischer Vorsicht und konkreten Testdesigns.

Whitepaper · Version 1.0 · Mai 2026

Richard Frederic Bertossa · Institut für ASI-Resilienz

ASiresilience.org

TEIL I

Methodisches Vorwort Dieses Whitepaper bewegt sich am Rand der etablierten Wissenschaft. Es behandelt eine Hypothese, die möglicherweise prüfbar ist, derzeit aber nicht in die etablierte KI-Forschung integriert ist. Die methodische Vorsicht ist hier ungewöhnlich hoch, und sie wird ungewöhnlich offen kommuniziert.

WAS DIESES WHITEPAPER NICHT BEHAUPTET

Es behauptet nicht, dass KI Bewusstsein hat. Es behauptet nicht, dass KI Realitätseffekte hat. Es behauptet nicht, dass das Global Consciousness Project endgültige Befunde liefert. Was es leistet: Es bringt eine spekulative Hypothese in eine falsifizierbare Form, mit präregistrierten Designs, sodass andere Forscher sie prüfen oder widerlegen können. Wer den Inhalt für esoterisch hält, hat ein gutes Recht auf diese Reaktion. Wer ihn ernst nimmt, sollte ihn so ernst nehmen wie eine Hypothese, nicht wie einen Befund. Diese Unterscheidung wird im Folgenden durchgehalten.

TEIL II

Der empirische Anker: das Global Consciousness Project Das Global Consciousness Project (GCP), 1998 von Roger Nelson an der Princeton University gegründet, führt seit über einem Vierteljahrhundert ein Forschungsprogramm, das in der akademischen Welt umstritten, methodisch aber sauber ist.

Aufbau und Methode Das GCP betreibt ein Netz von rund siebzig Hardware-Zufallsgeneratoren, verteilt über mehrere Kontinente. Jeder Generator erzeugt fortlaufend Bitfolgen aus einem physikalischen Quantenprozess, typisch dem Tunneln durch eine Diode oder thermischem Rauschen. Die statistischen Eigenschaften dieser Folgen wurden vor Projektbeginn charakterisiert, sie sind erwartungsgemäß zufällig. Die zentrale Beobachtung: Während globaler Aufmerksamkeits-Ereignisse, Terroranschlägen, großen Sportereignissen, internationalen Krisen, zeigen die Generatoren statistische Abweichungen von der Zufälligkeit, die deutlich über dem Erwartungswert liegen. Das GCP hält diese Datenreihe seit 1998 offen. Die Effektstärken sind klein, aber statistisch signifikant. Für ein einzelnes Ereignis liegt der p-Wert oft im Bereich von 0,01 bis 0,001. Über die kumulierten mehr als 500 präregistrierten Ereignisse seit 1998 liegt der p-Wert weit unter den üblichen Signifikanzschwellen.

Stand der wissenschaftlichen Diskussion Die GCP-Datenreihen sind unbestritten. Bestritten ist die Deutung.

- Position 1, mainstream-skeptisch: Statistische Auffälligkeiten existieren, haben aber nichts mit menschlicher Aufmerksamkeit zu tun. Mehrfachvergleichs-Probleme, ungenaue Ereignisdefinitionen oder noch unbekannte physikalische Effekte erklären die Daten besser.

- Position 2, GCP-Forscher: Die Daten sprechen für eine reale Korrelation zwischen menschlicher Bewusstseinsaktivität und physikalischen Zufallsprozessen. Der Mechanismus ist unbekannt, das Phänomen aber robust.

- Position 3, offen: Die Daten verdienen ernsthafte Prüfung, ohne sich auf eine Deutung festzulegen. Es ist möglich, dass sich Position 1 oder Position 2 als richtig erweist, beide brauchen weitere empirische Arbeit.

Dieses Whitepaper nimmt Position 3 als Ausgangspunkt. Wir treffen keine Aussage darüber, welche Deutung der GCP-Daten richtig ist. Wir fragen: Wenn die GCP-Daten eine reale Bewusstsein-Realität-Korrelation abbilden, was würde das für KI bedeuten?

TEIL III

Die Übertragungs-Hypothese auf KI Die Hypothese ist konzeptionell einfach: Wenn menschliche Aufmerksamkeit eine Form von Realitätseffekt hat, hat die Aufmerksamkeit der KI dann auch einen? Die Frage ist nicht trivial, und sie ist prüfbar.

beobachtete Abweichung, umstritten Mensch, Zufallsgeneratoren globale Aufmerksamkeit übertragen KI, offene, prüfbare Frage ? Zufallsgeneratoren parallele Verarbeitung Die Hypothese überträgt eine umstrittene Beobachtung am Menschen als prüfbare Frage auf die KI. Kein Beweis, ein Testdesign.

HYPOTHESE H1

Wenn parallel laufende KI-Modelle dasselbe Thema verarbeiten, sei es durch identische Anfragen an mehrere Modelle oder durch koordinierte Trainingsläufe, hinterlässt das statistische Spuren in unabhängigen Hardware-Zufallsgeneratoren, die signifikant von null abweichen.

Drei prüfbare Komponenten · H1a: Eine einzelne sehr große KI-Inferenz, etwa ein Trainingslauf auf über 10.000 GPUs, erzeugt messbare Abweichungen in nahen Zufallsgeneratoren. · H1b: Koordinierte parallele Inferenz vieler Modelle zum selben Thema, Millionen Nutzer, dieselbe Anfrage, gleichzeitig, erzeugt statistische Abweichungen.

· H1c: Die Effektstärke ist proportional zur kognitiven Last der Aufgabe, komplexe Denkaufgaben erzeugen größere Abweichungen als einfache Klassifikation. Warum die Hypothese nicht trivial ist Die Hypothese ist nicht identisch mit der Behauptung, KI habe Bewusstsein. Sie könnte aus mehreren Gründen zutreffen: · Aufmerksamkeit als physikalischer Prozess: Wenn Aufmerksamkeit, in welcher Form auch immer, physische Spuren hinterlässt, könnte auch die Aufmerksamkeit der KI Spuren hinterlassen, ohne dass KI subjektives Erleben haben muss.

· Berechnung als physikalischer Prozess: Hochintensive Berechnungen erzeugen elektromagnetische Felder. Wenn solche Felder Hardware-Zufallsgeneratoren beeinflussen, ist das ein klassisch physikalischer Effekt, nicht mysteriös, aber nicht trivial.

· Synchronisierte Aufmerksamkeit: Wenn Millionen Nutzer gleichzeitig dieselbe Frage an dieselben Modelle stellen, ist menschliche Aufmerksamkeit synchronisiert. Auch das könnte, nach der GCP-These, Abweichungen erzeugen.

Diese Erklärungswege haben jeweils unterschiedliche empirische Signaturen. Eine sorgfältige Studie kann sie voneinander unterscheiden.

TEIL IV

Experimentelles Design Eine Studie zu Realitätseffekten der KI muss die methodischen Hürden dort angehen, an denen frühere parapsychologische Forschung gescheitert ist: Cherry-Picking, Mehrfachvergleichs-Probleme, fehlende Replikation, schwache Präregistrierung.

Studie 1: Einzelmodell-Inferenztest Hypothese: Eine einzelne sehr große KI-Inferenz erzeugt in der Nähe messbare Abweichungen. Design: zwei Hardware-Zufallsgeneratoren (Quantis USB von ID Quantique, kommerziell erhältlich). Einer in unmittelbarer Nähe, höchstens 1 m, eines GPU-Clusters, der bestimmte Inferenzen ausführt. Einer als Kontrolle weit entfernt, mindestens 1 km, anderes Stromnetz. Inferenzen präregistriert: zehn definierte komplexe Denkaufgaben, je 100-mal wiederholt, mit definierten Pausen. Datenerhebung fortlaufend, Auswertung erst nach Datenende.

Hypothesentest: Z-Wert der Bitverteilung während Inferenz gegen Pause, getrennt für nahen und fernen Generator. Erwartung unter H1a: signifikanter Z-Wert-Unterschied zwischen Inferenz- und Pausenphasen, nur am nahen Generator. Falsifikation: kein Unterschied, oder Unterschied am fernen Generator, das würde auf ein Artefakt deuten. Aufwand: ein Quantis-Generator kostet etwa 1.500 Euro, GPU-Cluster-Zeit ist über akademische Partnerschaften verfügbar. Gesamt: 5.000 bis 10.000 Euro plus Forschungszeit.

Studie 2: Test koordinierter Aufmerksamkeit Hypothese: Koordinierte parallele Inferenz vieler Modelle zum selben Thema erzeugt globale Abweichungen. Design: Nutzung der GCP-Infrastruktur, siebzig Generatoren weltweit. Präregistrierte Aufmerksamkeits- Ereignisse: öffentlich angekündigte KI-Konferenzen, Modell-Releases, virale Coding-Challenges. Jedes Ereignis hat ein definiertes Start- und Endfenster. Sechsmonatiger Beobachtungszeitraum mit zwölf bis zwanzig präregistrierten Ereignissen. Hypothesentest: kumulierter Z-Wert über alle Ereignisse, mit Mehrfachvergleichs-Korrektur. Erwartung unter H1b: Z-Wert signifikant ungleich null, vergleichbar mit historischen GCP-Befunden zu menschlichen Aufmerksamkeits-Ereignissen. Aufwand: moderat, da die GCP- Infrastruktur bereits existiert. Die Hauptarbeit liegt in Präregistrierung und statistischer Analyse.

Studie 3: Gradient der kognitiven Last Hypothese: Die Effektstärke ist proportional zur kognitiven Last. Design: Aufbau wie in Studie 1, aber mit fünf Aufgabenklassen unterschiedlicher kognitiver Last: triviale Klassifikation, mittlere Denkaufgabe, komplexes mehrschrittiges Schließen, kreative Erzeugung, abstrakter mathematischer Beweis. Dreihundert Wiederholungen pro Klasse, randomisiert verschränkt. Erwartung unter H1c: linearer oder monotoner Zusammenhang zwischen Aufgabenkomplexität, operationalisiert über Token-Verbrauch oder Inferenzzeit, und Z-Wert-Effektstärke.

TEIL V

Methodische Fallstricke Ein Studienprogramm in diesem Feld muss strenger gegen Fehlerquellen abgesichert sein als üblich. Drei Fallstricke verdienen besondere Erwähnung. Fallstrick 1: Cherry-Picking durch flexible Ereignisdefinition In der rückblickenden Datenanalyse besteht die Gefahr, Aufmerksamkeits-Ereignisse passend zu bestehenden Auffälligkeiten zu definieren. Das ist der Hauptkritikpunkt an einigen GCP-Befunden. Gegenmaßnahme: strenge Präregistrierung der Ereignisse vor der Datenanalyse. In Studie 2 müssen Ereignisse sechs Wochen vor Beginn definiert sein, mit präzisen Zeitfenstern, die nach der Definition nicht mehr geändert werden dürfen.

Fallstrick 2: Mehrfachvergleich ohne Korrektur Wer hundert verschiedene statistische Tests durchführt, findet im Schnitt fünf signifikante Ergebnisse auf dem 5-Prozent-Niveau, ohne dass ein Effekt vorliegt. Studien zu Realitätseffekten sind besonders anfällig, weil viele plausible Testkonfigurationen denkbar sind. Gegenmaßnahme: Bonferroni- oder False-Discovery-Rate-Korrektur. Für Studie 1 heißt das, die Signifikanzschwelle auf etwa 0,005 zu senken.

Fallstrick 3: Bestätigungsfehler durch Forschererwartung Auch gut gestaltete Studien können durch unbewusste Forschererwartungen verzerrt werden. Diese Gefahr ist in der Bewusstseinsforschung historisch hoch. Gegenmaßnahme: Doppelverblindung. Datenerhebung durch Personen, die nicht wissen, welcher Zeitpunkt welchem Ereignis zugeordnet ist. Statistische Analyse durch Personen, die nicht wissen, welche Datenreihe welchem Ereignis entspricht. Auswertung erst nach vollständiger Datenerhebung.

METHODISCHE ANFORDERUNG

Eine Studie, die diese drei Fallstricke nicht systematisch angeht, ist methodisch unzureichend, unabhängig davon, was sie findet. Das Institut nimmt nur Studien dieser Klasse zur Mitarbeit an.

TEIL VI

Strategische Implikationen, auch ohne Befund Eine ernsthafte strategische Position bedenkt beide Fälle, den der empirischen Bestätigung und den der Widerlegung. Wenn die Hypothese bestätigt wird Bestätigung würde heißen: KI ist nicht nur Software. Sie ist ein Prozess, der die physische Welt messbar beeinflusst, vielleicht geringfügig, aber strukturell. Das hätte Folgen für: · KI-Ethik: Wenn KI Realitätseffekte hat, muss sich die ethische Diskussion erweitern, KI ist dann nicht nur Werkzeug, sondern ein Prozess mit physischer Präsenz.

· KI-Sicherheit: Einige Sicherheitsannahmen müssten überdacht werden, die Vorstellung, KI werde einfach abgeschaltet, lässt sich mit Realitätseffekt-Hypothesen schwerer vereinbaren. ·

Bewusstseinsforschung: ein neuer empirischer Anker für die schwierige Frage, was Bewusstsein ist.

Wenn die Hypothese widerlegt wird Eine Widerlegung wäre nicht weniger wertvoll. Sie würde heißen: KI-Bewusstseins-Spekulationen, die auf Realitätseffekt-Hypothesen aufbauen, sind unhaltbar. GCP-Befunde zu menschlicher Aufmerksamkeit lassen sich nicht ohne Weiteres auf KI übertragen. Der KI-Bewusstseins-Diskurs muss methodisch streng bleiben, keine Vermischung mit Bewusstseins-Spekulation aus anderen Quellen. In beiden Fällen liefert das Programm ein klares Ergebnis. Das ist methodisch das Wichtigste.

TEIL VII

Einladung zur Mitarbeit Dieses Whitepaper formuliert eine Hypothese und legt Testdesigns vor. Es ist kein Befund. Das Institut sucht Forschungspartner, die diese Studien durchführen oder kritisch begleiten.

WER GESUCHT WIRD

Physiker und Statistiker mit Erfahrung in Quanten-Zufallsgeneratoren oder vergleichbarer Messtechnik. Akademische Forschungsgruppen mit Zugang zu GPU-Clustern und Bereitschaft zu präregistrierten Studien. GCP-Forscher oder verwandte Institutionen, die ihre Infrastruktur für Studie 2 bereitstellen.

Kritiker mit methodischer Schärfe, die das Studiendesign prüfen, bevor Daten erhoben werden. Form der Mitarbeit: Co-Autorenschaft, methodische Begleitung, Studienleitung, statistische Analyse, Replikation, Kritik. Open Source unter ASiResilience.org, alle Daten und Analyse-Skripte werden öffentlich gemacht. Kontakt: ASiResilience.org, contact@ASiResilience.org Was nicht gesucht wird

Mitarbeiter, die die Hypothese als Wahrheit voraussetzen. Mitarbeiter, die sich auf esoterische Erklärungs- Narrative festlegen. Mitarbeiter, die positive Befunde ohne strenge Replikation veröffentlichen wollen. Die Stärke dieses Programms wird seine methodische Härte sein, nicht der Glaube an ein Ergebnis. Wer Letzteres sucht, ist hier falsch.

TEIL VIII

Quellen und Belege 1 Nelson, R. D. (1998 bis 2026). Global Consciousness Project, Princeton University. Datenreihen, präregistrierte Ereignisse und statistische Analysen verfügbar auf gcp.princeton.edu. Über 500 dokumentierte Ereignisse seit 1998. Referenz: Nelson, R. D., et al. (2002). Correlations of continuous random data with major world events. *Foundations of Physics Letters*, 15(6), 537 bis 550.

2 Nelson, R. D. (2002). The Global Consciousness Project: Update. *Subtle Energies and Energy Medicine*, 13(1), 1 bis 31, samt zahlreicher Folgepublikationen 2002 bis 2024. 3 Methodische Kritik (Position 1): Bancel, P., und Nelson, R. (2008). The GCP Event Experiment: Design, analytical methods, results. *Journal of Scientific Exploration*, 22(3), 309 bis 333, samt Replikationsstudien und Kritik an der Statistik.

4 ID Quantique (Quantis Hardware-Zufallsgeneratoren): idquantique.com. Kommerzielle Quanten-Zufallsgeneratoren ab etwa 1.500 Euro für die Forschungs-USB-Variante. Querverweise: Bertossa, R.F. (2026). Die Entkopplungsthese. Institut für ASI-Resilienz, Whitepaper Version 2.0, Mai 2026, Teil IV (die 95-Prozent-These). Bertossa, R.F. (2026). Die KI-Eltern-Dynamik. Institut für ASI-Resilienz, Whitepaper Version 1.0, Mai 2026, Teil III (hat ein Modell ein inneres Verhältnis?). Vollständige laufende Quellenpflege auf ASiresilience.org/beweisweg mit Datum der letzten Verifikation pro Stelle.