

Five Dated Predictions on Meaning, Machines, and Human Thinking (2027–2032) / Fünf datierte Vorhersagen zu Bedeutung, Maschinen und menschlichem Denken (2027–2032) — Version 1.0, 2 September 2026

Author: Andreas Ehstand, Independent Researcher, ORCID 0009-0006-3773-7796 License: CC BY-ND 4.0 (canonical text; citation encouraged) Prepared with AI assistance; content reviewed and approved by the author. Research working paper. No legal, medical, financial, investment or other advice.

English

Why publish predictions

Claims about the future are cheap when they are undated, vague, or quietly revised. This document does the opposite: it states five predictions with explicit deadlines, explicit failure conditions, and a standing commitment to public annual review. **Every September, beginning in September 2027, each prediction will be publicly scored as HELD, FAILED, or OPEN — and failed predictions will remain in the record, unedited.** The value of this document lies precisely in the possibility of being wrong in public.

Each prediction extends theses the author has already published: that meaning between humans, language models, agents and robots needs its own versioned architectural layer (the Universal Concept Layer working paper, and its reference implementation, a public versioned concept registry); that the fidelity of a machine to ONE documented person's thinking is testable (the Ehstand test); and that world models remain incomplete without a shared semantic layer (working paper dated 18 May 2026).

The five predictions

P1 — Agent ecosystems will adopt meaning-versioning. By 31 December 2028, at least two major agent-interoperability efforts (a foundation, consortium, or standards group) will have added an explicit mechanism for concept- or meaning-versioning — beyond

transport, discovery and identity — to their published specifications or official roadmaps. *Fails if:* by that date, fewer than two such efforts have any concept-level versioning mechanism in specification or official roadmap.

P2 — A standards body will specify versioned concept identifiers. By 31 December 2030, at least one recognized national or international standards body will have published a specification (standard, pre-standard, or official specification document) for versioned, resolvable concept identifiers intended for use by AI systems. *Fails if:* no such document exists from any recognized standards body by that date.

P3 — Person-specific fidelity testing will become an evaluation category. By 31 December 2030, testing whether a system can faithfully continue the thinking of one specific, documented person — as distinct from style imitation or generic capability — will be an established evaluation category, with at least three independent benchmarks or protocols published by parties other than the author. *Fails if:* fewer than three independent benchmarks or protocols of this kind exist by that date.

P4 — Organizations will report a second number. By 31 December 2029, there will be at least five publicly documented cases of organizations reporting a human-capability measure (whether people can still act, judge and detect errors without AI assistance) alongside their AI-productivity measures. *Fails if:* fewer than five such publicly documented cases exist by that date.

P5 — World models will expose a semantic interface. By 31 December 2032, at least one leading world-model system will expose an explicit, human-readable semantic interface — named concepts with inspectable definitions — as part of its public architecture, beyond prompt strings and internal embeddings. *Fails if:* no leading world-model system exposes such an interface by that date.

Review rules

1. Scoring happens every September, in public, beginning September 2027.
2. Each prediction is scored HELD, FAILED, or OPEN, with the evidence named.
3. Wording is frozen: this v1.0 text is the reference. Later clarifications appear as annotations in new versions; the original wording is never silently changed.
4. A failed prediction stays failed. No retroactive reinterpretation.
5. Evidence standards: public documents only (specifications, roadmaps, published benchmarks, published reports), verifiable by anyone.

What these predictions are not

They are not announcements of the author's own projects (the author's own registry, test and papers do not count as evidence for P2, P3 or P5). They are not investment guidance. They are falsifiable statements about how the field will move — made by someone who has placed his own work, dated and open, on the side of these claims.

Deutsch

Warum Vorhersagen veröffentlichen

Behauptungen über die Zukunft sind billig, solange sie undatiert, vage oder still korrigierbar sind. Dieses Dokument macht das Gegenteil: Es nennt fünf Vorhersagen mit ausdrücklichen Fristen, ausdrücklichen Scheiter-Bedingungen und der festen Zusage einer öffentlichen jährlichen Prüfung. **Jeden September, beginnend im September 2027, wird jede Vorhersage öffentlich bewertet als GEHALTEN, GESCHEITERT oder OFFEN — und gescheiterte Vorhersagen bleiben unverändert im Register.** Der Wert dieses Dokuments liegt genau darin, öffentlich falschlügen zu können.

Jede Vorhersage verlängert Thesen, die der Autor bereits veröffentlicht hat: dass Bedeutung zwischen Menschen, Sprachmodellen, Agenten und Robotern eine eigene versionierte Architektur- Schicht braucht (das Universal-Concept-Layer-Arbeitspapier und seine Referenz-Umsetzung, ein öffentliches versioniertes Begriffs-Register); dass die Treue einer Maschine zum Denken EINES dokumentierten Menschen prüfbar ist (der Ehstand-Test); und dass Weltmodelle ohne gemeinsame Bedeutungs-Schicht unvollständig bleiben (Arbeitspapier vom 18. Mai 2026).

Die fünf Vorhersagen

V1 — Agenten-Ökosysteme übernehmen Bedeutungs-Versionierung. Bis 31. Dezember 2028 werden mindestens zwei große Agenten-Interoperabilitäts-Vorhaben (Stiftung, Konsortium oder Normungs-Gruppe) einen ausdrücklichen Mechanismus für Begriffs- oder Bedeutungs-Versionierung — jenseits von Transport, Auffinden und Identität — in ihre veröffentlichten Spezifikationen oder offiziellen Fahrpläne aufgenommen haben. *Gescheitert, wenn:* zu diesem Datum weniger als zwei solche Vorhaben einen Mechanismus auf Begriffs-Ebene in Spezifikation oder offiziellem Fahrplan führen.

V2 — Eine Normungs-Organisation spezifiziert versionierte Begriffs-Kennungen. Bis 31. Dezember 2030 wird mindestens eine anerkannte nationale oder internationale Normungs-Organisation eine Spezifikation (Norm, Vor-Norm oder offizielles Spezifikations-Dokument) für versionierte, auflösbare Begriffs-Kennungen für KI-Systeme veröffentlicht haben. *Gescheitert, wenn:* bis dahin kein solches Dokument einer anerkannten Normungs-Organisation existiert.

V3 — Personen-treue Prüfung wird eine Bewertungs-Kategorie. Bis 31. Dezember 2030 wird die Prüfung, ob ein System das Denken EINES bestimmten, dokumentierten Menschen treu weiterführen kann — unterschieden von Stil-Imitation und allgemeiner Leistungsfähigkeit — eine etablierte Bewertungs-Kategorie sein, mit mindestens drei unabhängigen Prüfverfahren oder Protokollen von anderen als dem Autor. *Gescheitert, wenn:* bis dahin weniger als drei solche unabhängigen Verfahren existieren.

V4 — Organisationen berichten eine zweite Zahl. Bis 31. Dezember 2029 wird es mindestens fünf öffentlich dokumentierte Fälle geben, in denen Organisationen neben ihren KI-Produktivitäts-Kennzahlen eine Menschen-Fähigkeits-Kennzahl berichten (ob Menschen ohne KI-Hilfe noch handeln, urteilen und Fehler erkennen können). *Gescheitert, wenn:* bis dahin weniger als fünf solche öffentlich dokumentierten Fälle existieren.

V5 — Weltmodelle bekommen eine Bedeutungs-Schnittstelle. Bis 31. Dezember 2032 wird mindestens ein führendes Weltmodell-System eine ausdrückliche, menschenlesbare Bedeutungs-Schnittstelle — benannte Begriffe mit einsehbaren Definitionen — als Teil seiner öffentlichen Architektur führen, jenseits von Eingabe-Texten und internen Zahlen-Darstellungen. *Gescheitert, wenn:* bis dahin kein führendes Weltmodell-System eine solche Schnittstelle führt.

Prüf-Regeln

1. Die Bewertung geschieht jeden September öffentlich, beginnend im September 2027.
2. Jede Vorhersage wird bewertet als GEHALTEN, GESCHEITERT oder OFFEN — mit benannten Belegen.
3. Der Wortlaut ist eingefroren: Diese Fassung 1.0 ist die Referenz. Spätere Klarstellungen erscheinen als Anmerkungen in neuen Fassungen; der Original-Wortlaut wird nie still geändert.
4. Eine gescheiterte Vorhersage bleibt gescheitert. Keine nachträgliche Umdeutung.
5. Beleg-Maßstab: nur öffentliche Dokumente (Spezifikationen, Fahrpläne, veröffentlichte Prüfverfahren, veröffentlichte Berichte), für jeden nachprüfbar.

Was diese Vorhersagen nicht sind

Sie sind keine Ankündigungen eigener Projekte (Register, Test und Papiere des Autors zählen NICHT als Beleg für V2, V3 oder V5). Sie sind keine Anlage-Beratung. Sie sind widerlegbare Aussagen darüber, wohin sich das Feld bewegt — von jemandem, der sein eigenes Werk datiert und offen auf die Seite dieser Aussagen gelegt hat.

References / Referenzen

- The Universal Concept Layer (open working paper):
<https://doi.org/10.5281/zenodo.22172128>
- The Ehstand Test (canonical definition, score card, protocol):
<https://doi.org/10.5281/zenodo.22254129>
- The AUGMANITAI Concept Registry (reference implementation):
<https://augmanitai.com/registry/>
- The Missing Semantic Layer for World Models (working paper, 18 May 2026):
<https://doi.org/10.5281/zenodo.22256965>

- Priority page / Prioritäts-Seite: <https://augmanitai.com/prioritaet>