

The Attralucian Essays:

Exploring the Finite



First Edition

Copyright © 2025 by Kevin R. Haylett. All rights reserved.

This work is shared under the Creative Commons Licence.

Creative Commons CC BY-ND 4.0 License.

<https://creativecommons.org/licenses/by-nd/4.0/>

This work is intended for academic and research use. Any unauthorized distribution, modification, or commercial use beyond the creative use license is strictly prohibited. Typeset in

L^AT_EX

The Attralucian Essays



Semantic Uncertainty: Towards Semantic
Accountability in Scientific Discourse

Kevin R. Haylett

Semantic Uncertainty: Towards Semantic Accountability in Scientific Discourse

Kevin R. Haylett

Astract

Scientific language, particularly in theoretical and interdisciplinary domains, frequently relies on abstract terms such as “consciousness,” “intelligence,” or “information” without disclosing the underlying semantic scaffolding upon which these terms rest. While numerical models often include uncertainty quantification, linguistic formulations are typically presented as semantically stable despite operating in inherently unstable meaning-spaces. This paper proposes a structured system of semantic accountability through the addition of a Semantic Uncertainty Appendix in all theory-heavy or language-dependent research outputs. Drawing on the ”words as

transducers” framework introduced in Attralucian Essays (01), we argue that words function analogously to sensors—measuring, compressing, and projecting both internal and external structures. As such, they inherently carry measurement uncertainty. We offer a rationale, formal structure, and illustrative examples to support this addition to scientific practice.

Introduction

In quantitative science, it is standard practice to disclose the uncertainty bounds of any numerical measurement. A voltage reading of 2.21 V might be annotated as ± 0.01 V, reflecting known variability in instrumentation or environment. However, in theoretical or interdisciplinary writing, we frequently encounter claims such as “consciousness arises from microtubule coherence” or “language models exhibit understanding” without any corresponding notation of semantic variability or scope. Yet these terms—“consciousness,” “understanding,” “free will,” “representation”—are not fixed. They operate as semantic attractors with highly variable boundaries, histories, and internal contradictions. Their usage without clarification introduces latent instability into the discourse, often recognized as theoretical disagreement rather than linguistic drift.

This essay proposes that science adopt a systematic approach to semantic uncertainty, just as it does for numer-

ical and methodological uncertainty. Such an approach can be normalized through a structured appendix included with relevant publications, outlining: Operational definitions of key terms;

Known ambiguities or recursive risks;

Acknowledged metaphors and analogies;

Domains of valid application and drift;

Justification for terminological choices

The Transducer model of language

In *Finite Models of Words: Words as Transducers* (Haylett, 2025), words are modeled not as static symbols but as transducers—finite operators that compress, project, and translate between structures, both internal (e.g., latent semantic geometries) and external (e.g., measured physical data). Each word, in this view, performs semantic compression: “Fire” encodes a learned manifold of heat, color, hazard, emotion, and causality; “Consciousness” compresses first-person awareness, functional cognition, philosophical dualisms, and cultural intuitions; “Free will” collapses centuries of metaphysical debate into a grammatical noun phrase. Just as with physical sensors, these transductions are finite, lossy, and context-sensitive. They must therefore be treated as measurements with embedded uncertainties.

0.0.1 The Proposal

We propose the formal adoption of a Semantic Uncertainty Appendix (SUA) for all research that makes significant use of theoretical or cross-domain terms whose meanings are known to shift.

Objectives:

Improve semantic transparency;

Reduce disciplinary miscommunication;

Enhance AI interpretability and training quality;

Expose hidden assumptions behind theoretical claims;

Enable recursive reflection for future work.

Term	Operational Definition	Known Ambiguities	Context of Validity	Justification
Consciousness	The unified subjective experience modeled as collapse events (Orch OR).	Conflates functional, phenomenological, and cultural meanings. Risks recursion (“awareness of awareness”).	Neurocomputational models. Excludes cultural or mystical frames.	Chosen for alignment with cited empirical studies.
Intelligence	The capacity to navigate semantic attractors toward goal-directed inference.	Often implies agency or generalization not modeled.	Applicable within LLM phase space interpretation.	Preferable to “understanding”; introduces recursive instability.
Binding	Phenomenological unification of sensory or representational elements.	Used in neural, linguistic, and philosophical senses. May assume unity where none exists.	Cognitive neuroscience and representational linguistics.	Retained for its utility across interpretive frames.

Table 1: Recommended Operational Definitions and Contexts (Landscape View)

Practical Concerns

Resistance to implementation will likely be rooted in: Perceived “softness” of semantics vs. numerical results; Fear of exposing internal uncertainty as weakness; Lack of training in meta-linguistic reflection. We argue, however, that these are the same barriers once faced by uncertainty quantification in physics, biology, and engineering. Normalizing semantic reflection will similarly improve precision, falsifiability, and epistemic humility

0.0.2 Enhancing Public Trust and Science Communication

The adoption of SUAs could significantly improve public engagement with science by demystifying the linguistic foundations of complex claims. In an era where terms like “artificial intelligence” or “quantum consciousness” are often sensationalized in popular media, an SUA provides a transparent framework for clarifying what scientists mean—and what they don’t. By openly acknowledging semantic uncertainties, researchers can preempt misinterpretations that fuel public skepticism or distrust, particularly in contentious fields like AI ethics or climate science. This transparency aligns with broader efforts to make science more accessible and accountable, fostering a culture where uncertainty is seen not as weakness but as a hallmark of rigorous inquiry. Furthermore, SUAs could serve as a bridge between scientific discourse and public

communication. Science communicators and journalists could draw on these appendices to craft narratives that accurately reflect the scope and limitations of research, reducing the risk of hype or distortion. For instance, an SUA clarifying the term “sustainability” in environmental studies could guide reporters to avoid overly broad or misleading claims. By embedding semantic accountability into the research process, the SUA empowers scientists to shape public discourse proactively, ensuring that their work is both understood and trusted.

Conclusion

Words are not neutral. They are finite instruments of measurement—transducers—that inherit the variability and drift of their contexts. To treat them as fixed is to build conceptual cathedrals on unstable ground. Just as we would not publish a voltage measurement without a margin of error, we must not publish abstract theoretical claims without a transparent framing of semantic uncertainty. In an age where language is both the instrument and object of artificial cognition, semantic accountability is not optional. It is foundational