

LUIS ALBERTO MONTOYA-PEREZ

Research Lineage

Before the models, the pattern was already there.

Selected historical academic work and methodological development · 2014 → now

HISTORICAL COURSEWORK

DEVELOPMENTAL PROVENANCE

NOT CURRENT VALIDATION

The public thesis

The archive documents continuity of research habits before contemporary generative AI: source checking, contextual interpretation, observation, systems reconstruction, explicit critical-thinking frameworks, causal restraint, and later quantitative exposure. It also preserves mistakes and overreach. The modern controls are presented as later formalizations, not retroactive properties of the student work.

The lineage at a glance

Period	Development	Evidence	Thread carried forward
2014	Source checking	Official source + independent source	SINGLE SOURCE → CROSS-CHECKING
2016	Naturalistic observation	Preschool field observation	OBSERVATION → INFERENCE → ALTERNATIVES
2017	Systems reconstruction	Spectroscopy as cumulative lineage	PARTIAL FINDINGS → LINEAGE → SYSTEM
2018	Verification + interrupted direction	Myth / source verification; psychology-oriented return to school later withdrawn	CLAIM → SOURCE → VERIFICATION
2022	Explicit critical thinking	English 1C under Prof. M. Hanson	IMPLICIT BRIDGE → EXPLICIT REASONING STRUCTURE
2022	Scientific restraint	BIOL2 measurement, causation, bias, ethics	PLAUSIBLE EXPLANATION → BOUNDED CLAIM
Later	Formal statistics	Study design,	QUANTITATIVE

		probability, hypothesis testing	EXPOSURE → SPECIALIST BOUNDARY
2026	Externalized method	Provenance, claim ceilings, adversarial review	CAPABILITY + CONTROL

2018 is an interrupted chapter, not a flagship credential

Retrospective context: in 2018 I returned to school still oriented toward psychology, but without a clear path for how to move forward. Overtime work and school competed for the same time, and I withdrew when the load became unsustainable. That period belongs in the lineage because it documents uncertainty and interruption rather than a clean upward academic story.

Provenance boundary

A surviving 2018 Freud/Jung/Lacan paper is labeled English 1C · Module 7 in the document itself. I am holding it out of the public flagship set because the administrative sequence of multiple English 1C attempts is not fully reconstructed, and my recollection distinguishes the later Hanson experience as the developmentally important one.

2018 · Myth, verification, and persuasion

A separate 2018 course submission explicitly discusses content verification, source verification, propaganda, critical thinking, and emotional appeal. Its importance is historical orientation, not factual authority.

2022 · English 1C makes the reasoning structure explicit

The surviving English 1C folder contains about 57 rendered pages of distinct substantive coursework after removing one text-identical copy and treating one near-duplicate essay pair as versions of the same workstream. The archive spans more than one attempt, so individual files are not assigned to a failed versus withdrawn attempt without stronger administrative evidence.

The strongest artifact is the 03 October 2022 Module 7 Review. It contemporaneously records training in context and subtleties, visual connection of related ideas, critical thinking, intellectual standards, elements of reasoning, propaganda, rhetorical dimensions, supporting evidence, revision, fact/opinion separation, and deconstructing an idea.

What this establishes

The course explicitly taught practices strongly aligned with later conceptual and methodological interests. It does not establish that Professor Hanson caused the present methodology, nor that today's methods existed unchanged in 2022.

2022 • BIOL2 exposes both the strength and the failure mode

- **Operationalization:** The giraffe-study entry focuses on how food consumption was measured, including bite counts and dry mass per bite.
- **Evidence-quality ceiling:** The same entry says low image resolution prevented confident evaluation of the result.
- **Causal restraint:** A dog-behavior genetics entry explicitly states that correlation does not imply causation.
- **Agency and ethics:** The genetic-testing entry weighs utilitarian and libertarian perspectives and recognizes rights both to know and not to know.
- **Bias:** A biodiversity entry identifies taxonomic bias and argues that broader parameter coverage can mitigate, but not eliminate, research bias.
- **Source-access boundary:** A later entry explicitly states that a paywalled source could not be read, then unfortunately makes an unsupported downstream claim. Both the good boundary and the overreach are preserved.

The developmental finding

The strength Across unrelated domains, the record repeatedly shows rapid construction of relationships between sources, contexts, observations, theories, and systems.
The failure mode Some historical bridges are more elegant than the evidence can support. Construct mismatch, confounding, causal compression, and over-transfer appear in the record.
The later control The modern emphasis on provenance, claim ceilings, specialist authority, competing explanations, falsifiability, and adversarial review is a rational control around a documented cognitive strength and its documented failure mode.

From historical pattern to current method

Historical pattern	Current methodological form
Cross-check sources	provenance and source authority
Context changes interpretation	context-shift / frame-fidelity evaluation
Observation is not explanation	inference ceilings / competing explanations
Decompose reasoning	separate evaluation dimensions / anti-masking
Correlation $\neq$ causation	design-dependent causal claims
Research carries bias	investigator/process threat registers
Bridges can outrun evidence	adversarial review / transfer limits
Quantitative knowledge has limits	qualified statistical ownership

Current research • AI for Wellbeing

The current proposal invites The Trevor Project, Anthropic, and OpenAI to help test which small, controlled changes in context or AI response matter downstream, under what conditions, and with what evidence for causation.

TRACK A • AI BEHAVIOR	TRACK B • HUMAN OUTCOMES
Tests recalibration as evidence and context change, with separate scoring for frame reconstruction, state-transition calibration, factual handling, agency, harmful agreement, and unnecessary refusal.	Studies bounded response differences and human outcomes, with a proposed adolescent focus subject to matching expertise, ethics review, safeguarding, consent/assent, funding, and final design.
Independent adversarial review The exact AI for Wellbeing v0.4 proposal received a SHA-bound external PASS for methodological coherence and readiness for external partner review. SHA-256: 8dd8d966f20bfa80906a9d581421b270153afa70a006c1c5a56ac02476eb732c This is not empirical validation, partner endorsement, clinical authority, human-subjects approval, or evidence that any proposed organization has agreed to participate.	

Public artifact set

Artifact	Status	Reason
2014 Google Glass summary	Public	Earliest reviewed source-cross-checking anchor.
2016 preschool observation	Public, redacted	Primary observational provenance; identities removed.
2017 spectroscopy	Public	Clean systems/lineage reconstruction.
2018 myth/verification essay	Public archival edition	Source verification, propaganda, critical-thinking lineage.
2022 English 1C Module 7 Review	Public archival edition	Contemporaneous record of explicit reasoning instruction.
2022 BIOL2 method excerpts	Public curated edition	Measurement, causation, ethics, bias, source-access boundaries.
2022 Sun vs. Shade observation	Public archival edition	Observation → candidate hypothesis.
2018 Freud/Jung/Lacan	Held out	Course-attempt attribution unresolved; not needed for the public thesis.
2023 statistics assessments	Private/supporting	Formal exposure with visible limits; not a public accuracy credential.

What this archive does not claim

- Historical coursework qualifies me to provide clinical or crisis intervention.

- Old papers validate the current AI for Wellbeing, DVS, Longitudinal, ORPHARD, Persona, RVEA, or other modern methods.
- The English 1C sequence was completed in a single clean attempt.
- The 2018 Freud paper is the definitive English 1C flagship.
- The surviving 57-page corpus was definitely submitted as one 57-page final portfolio.
- Formal statistics coursework makes me the statistical authority for a human-outcomes study.
- The historical work was uniformly factually correct.
- Trevor, Anthropic, OpenAI, Daniel J. Puhlman, University of Maine, or Professor Hanson endorses the work unless and until that is separately established.

Identity and verification

Luis Alberto Montoya-Perez · Originating investigator and evaluation architect at unFragged.

Selected academic archive: Barstow Community College coursework and related course exports.

Current interests: AI evaluation, context-sensitive response behavior, evidence provenance, long-horizon human/AI interaction, and wellbeing research.