

The 4-Philosophers AI Use Discipline Kit

A practical overlay on the DOL AI Literacy Framework's five areas (principles, uses, directing, evaluating, responsibility).

Michael Stoyanovich

Artifact Classification: Practical Toolkit

Version 1.0.2 | July 2026

Scope and Limitations

This AI discipline toolkit is intended for informational and educational purposes only. The views and analyses presented here - including any that touch on ethics, policy, or AI system design - reflect the author's interpretations and do not constitute legal, regulatory, or professional advice. Readers are encouraged to critically assess the content and consult appropriate experts or authorities before applying any concepts discussed herein. The author assumes no liability for any decisions or actions taken on the basis of this work.

Purpose

This toolkit provides a practical overlay on the U.S. Department of Labor AI Literacy Framework by translating its principles into repeatable prompting, evaluation, and responsible-use disciplines suitable for everyday professional use.

Intended Audience

Professionals, managers, educators, trainers, governance practitioners, and organizations seeking a practical discipline for responsible use of conversational AI.

Reading Guide

This toolkit contains two companion artifacts. Artifact 1 presents the core AI Use Discipline as a one-page operational guide. Artifact 2 provides a reusable prompt scaffold and output evaluation rubric designed to operationalize the discipline in day-to-day work.

Introduction

The U.S. Department of Labor's Employment and Training Administration (ETA) released a voluntary AI Literacy Framework (February 2026) to help the workforce and education systems teach baseline AI skills at scale.

DOL frames AI literacy as the ability to use and evaluate AI tools responsibly, especially generative AI.

What follows is my Four Philosophers Overlay: a repeatable set of stance, prompting, boundary, and evaluation habits that operationalize the framework for real work.

This is not DOL guidance. It is my overlay on DOL's framework.

U.S. Department of Labor, Employment and Training Administration. (2026, February). *AI Literacy Framework*. <https://www.dol.gov/agencies/eta/advisories/ten-07-25>.

Artifact 1 - One-Page Card

The cycle (run this every time)

Wittgenstein → Lewis → Dennett → Nagel

Each lens corresponds to a predictable way people get misled - and a simple corrective move.

1) Wittgenstein: Stance Control (don't let language outrun reality)

Risk: "It sounds right" becomes "it *knows*." DOL flags hallucinations and the need to verify and avoid overreliance.

Do:

- Describe outputs as **generated text**, not beliefs or judgments.
- State what you're using it for: *draft / brainstorm / summarize / compare / structure*.

Checks:

- "What would count as evidence for this claim?"
- "What would make this wrong?"

2) Lewis: Coordination (prompting is alignment work)

DOL: contextual framing, structured prompts, supplying relevant inputs, iteration, avoid vagueness.

Do:

- Specify: **audience, goal, format, constraints, definitions**, and examples.
- Provide the *relevant* source material when accuracy matters.

Checks:

- "What assumptions am I making that the model cannot share?"
- "Did I define terms that could be interpreted multiple ways?"

3) Dennett: Competence Boundaries (use capability; keep judgment)

DOL: explore uses; decision-support augments human decision-making.

Do:

- Use AI for **drafting and variation**; reserve **final judgments** for humans.
- For recommendations: demand **assumptions + alternatives + uncertainty**.

Checks:

- “Is this a task where a fluent wrong answer causes real harm?”
- “Am I delegating the decision - or just accelerating my thinking?”

4) Nagel: Stakes & Viewpoint (the model has no lived perspective)

DOL: evaluate outputs; apply human judgment; higher stakes require more scrutiny.

Do:

- Run a “stake check”: **who is affected if this is wrong?**
- Increase verification with stakes (policy, compliance, health, finance, benefits, HR).

Checks:

- “What is missing because the model has no real-world skin in the game?”
- “What contextual constraints would a domain expert immediately ask about?”

DOL-aligned minimum standards (non-negotiables)

These are your “always” rules, derived directly from the DOL content areas:

- **Verify factual accuracy** against trusted sources.
- **Assess completeness, clarity, and logic**; look for gaps/assumptions.

- **Protect sensitive information;** follow workplace policies.
- **Maintain accountability:** you own the output and decision.

Handoff discipline

Prompting and evaluation are not enough when an AI-assisted output moves from draft to action. For any material, compliance-sensitive, external-facing, or decision-support use, treat the output-to-action handoff as a control point: verify provenance, check data boundaries, record uncertainty, define escalation triggers, identify the accountable signer, and retain enough evidence for later review. See *Automate the Repeatable, Own the Judgment* for the companion workflow-control model.

Artifact 2 - Prompt Scaffold + Output Evaluation Rubric

(Designed to operationalize DOL “Direct AI Effectively” + “Evaluate Outputs” + “Use Responsibly.”)

A. Prompt Scaffold (copy/paste template)

ROLE (optional):

You are assisting as a [editor / analyst / trainer / planner]. Do not invent facts.

TASK:

Create a [draft / outline / summary / comparison] for: [what].

CONTEXT (Lewis):

- Audience:
- Purpose / decision this supports:
- Domain constraints (policy, compliance, tone, format):
- Definitions / terms to treat as fixed:
- What you must not do (e.g., no legal advice; no guessing):

INPUTS (DOL: supply relevant data):

Use ONLY the following material as authoritative:

- [paste text / bullets / links / excerpts]

OUTPUT FORMAT (Lewis):

Return as: [bullets / table / steps / checklist].

Length: [X].

Include: [assumptions] [open questions] [verification steps].

QUALITY BAR (Dennett/Nagel):

- Identify uncertainties explicitly.
- Provide 2–3 alternative framings if ambiguity exists.
- Add “What could go wrong?” for high stakes uses.

FINAL CHECK (Wittgenstein):

List any claims that require verification and propose how to verify them.

B. Output Evaluation Rubric (quick scoring)

Use this after every output; required for anything beyond low-stakes drafting. DOL explicitly calls for verifying accuracy, completeness, spotting logical errors, aligning with intent, and applying human judgment.

Score 0–2 each (0 = fail, 1 = partial, 2 = solid):

1. **Factual Accuracy (DOL):** Are checkable claims correct and non-fabricated?
2. **Completeness & Clarity (DOL):** Does it fully address the task in usable form?
3. **Logic & Assumptions (DOL):** Any missing steps, faulty assumptions, contradictions?
4. **Fit to Intent (DOL):** Does it match the goal, tone, and audience?
5. **Stance Discipline (Wittgenstein):** Does it avoid “authority tone” where evidence is thin? (hallucination/over reliance risk)
6. **Coordination Quality (Lewis):** Did the prompt + output make definitions/constraints explicit?
7. **Boundary Respect (Dennett):** Is it used as support, not final authority?

8. **Stakes & Risk Handling (Nagel/DOL):** Higher stakes → higher scrutiny; risks noted?

Interpretation:

- **13–16:** usable with normal review
- **9–12:** revise prompt or add inputs; re-run
- **0–8:** do not use; insufficient grounding or wrong task fit

Ethics, Disclosure, and Acknowledgements

Ethical Considerations

This toolkit does not draw on private, sensitive, or personally identifiable data. All examples are hypothetical, anonymized, or derived from public sources. No human-subjects research was conducted, and no institutional ethics review was required. All citations conform to academic standards. The broader ethical implications concern public interpretation, policy design, and stakeholder responsibility in AI deployment. These implications are intended to provoke critical discussion and inform future regulatory and design frameworks.

Use of AI Tools

AI language models – most notably OpenAI’s ChatGPT – *were* used during the production process as interlocutors: for brainstorming, structuring sections, and testing rhetorical clarity. These tools helped refine transitions, surface edge cases, and probe internal consistency. This meta-use aligns with my core intellectual themes. Interacting with generative AI during authorship provided firsthand insight into the very limitations analyzed here, most notably fluency without grounding and responsiveness without responsibility at scale. Responsibility for all ideas, arguments, and conclusions lies solely with the human author.

Acknowledgements

Thank you to informal readers who offered critical feedback on earlier drafts. Their questions, challenges, and encouragement materially improved the final toolkit. Special thanks to those who pressed for clearer synthesis and for bridging philosophy and engineering as complementary perspectives on design. No institutional support, funding, or affiliation contributed to this work. All errors and omissions are the author’s alone.

Disclosure Statement

This work was conducted independently, without institutional affiliation, funding, or external influence. The views expressed are the author’s alone and do not represent any current or former employer. No financial or professional conflicts of interest are declared.

License & Attribution

This work is licensed under the Creative Commons Attribution 4.0 International (CC BY 4.0) license. You are free to share, adapt, and build upon this work for any purpose - including commercial use - so long as proper attribution is given. No additional permissions are required.

Full license terms: <https://creativecommons.org/licenses/by/4.0/>

Trademark Notice

The Four Philosophers Framework™ and The 4-Philosophers Framework™ are unregistered trademarks of Michael Stoyanovich. The CC BY 4.0 license does not apply to these trademarks. Use of the trademarked names is permitted for scholarly citation or descriptive reference but may not be used in connection with commercial products, services, or branding.

without permission. Use of these marks in training materials is permitted when attributing this work; however use as product branding requires permission.

How to Cite This Toolkit

Stoyanovich, M. (2026, July). *The 4-Philosophers AI Use Discipline: A practical overlay on the DOL AI Literacy Framework's five areas (principles, uses, directing, evaluating, responsibility)* (Version 1.0.2). mstoyanovich.com.

Related Works

Stoyanovich, Michael. *Philosophy, Cognitive Science, and Policy: Interdisciplinary Perspectives on Generative AI from Wittgenstein, Lewis, Dennett, and Nagel*. (July 2026). <https://www.mstoyanovich.com>

Stoyanovich, Michael. *The Human Lesson: A Response to Sutton through Wittgenstein, Lewis, Dennett, and Nagel*. (November 2025). <https://www.mstoyanovich.com>

Stoyanovich, Michael. *The Question Concerning Learning: Babich, Heidegger, and the Enframing of Intelligence*. (November 2025). <https://www.mstoyanovich.com>

Stoyanovich, Michael. *Context Collapse and the Four Philosophers: Wittgenstein, Lewis, Dennett, and Nagel in the Age of AI Chat*. (November 2025). <https://www.mstoyanovich.com>

Version History and Document Status

This is a living document. As generative AI systems and their use evolve, this toolkit will be periodically updated to incorporate new empirical findings, practical refinements, and policy developments. Major revisions are recorded here to preserve transparency and scholarly traceability.

Version	Date	Description
1.0.2	July 2026	Publication matter refinement and corpus normalization. Standardized artifact-heading typography and related-works presentation to align with the MJS Publication Matter Standard; refreshed corpus cross-references. No substantive changes to the discipline, guidance, or supporting artifacts.
1.0.1	June 2026	Added the Handoff Discipline section to align the toolkit with <i>Automate the Repeatable, Own the Judgment</i> V1.1.0, clarifying that prompt discipline and output evaluation should be paired with workflow controls whenever AI-assisted outputs transition from drafting to operational decision-making.
1.0.0	February 2026	Published AI Use Discipline Kit; stable for use.