

BEYOND HALLUCINATION: Semantic Laundering, Semantic Stewardship, and Prompting Aikido in Human-AI Collaboration

Dennis Kennedy, Kennedy Idea Propulsion Laboratory (KIPL)

Suggested Citation: Kennedy, D. (2026). Beyond Hallucination: Semantic Laundering, Semantic Stewardship, and Prompting Aikido in Human-AI Collaboration. SSRN Electronic Journal.

License: Creative Commons Attribution 4.0 International (CC BY 4.0)

Author: Dennis Kennedy*

Affiliation: Kennedy Idea Propulsion Laboratory (KIPL)

Keywords: Generative AI, Semantic Laundering, Semantic Stewardship, Prompting Aikido, Aikido Prompting, Human-AI Collaboration, Knowledge Work

JEL Classifications: O33, K40, D83

** **AI Disclosure Statement:** Generative AI tools (including Claude and Gemini) were utilized during the preparation of this work as collaborative co-reasoning partners for structural feedback, preliminary drafting, and formatting. In accordance with the semantic stewardship methodology detailed herein, all concepts, arguments, and final text were critically evaluated, substantively revised, and approved by the author, who assumes full responsibility for the manuscript.*

- **Author:** Dennis Kennedy
- **Affiliation:** Kennedy Idea Propulsion Laboratory (KIPL) / Michigan State University College of Law
- **Keywords:** Generative AI, Semantic Laundering, Semantic Stewardship, Prompting Aikido, Human-AI Collaboration, Knowledge Work, Legal Tech
- **JEL Classifications:** O33 (Technological Change), K40 (Legal System: General), D83 (Search, Learning, Information and Knowledge)

Abstract

The primary conversation surrounding generative artificial intelligence centers on factual fabrication: hallucination. While hallucination matters, it represents a crude and easily detected failure mode. This paper identifies a second, quieter epistemic risk: semantic laundering. Large language models produce exceptionally fluent, coherent, and orderly prose that silently smooths away anomalies, unresolved tensions, weak signals, idiosyncratic language, and vital distinctions. In high-stakes domain inquiry, this polite normalization erodes the very irregular material on which professional judgment depends. Drawing on a reflexive design case, this paper proposes semantic stewardship as a governing practice for preserving meaningful differentiation. It introduces prompting aikido: compact, context-sensitive micro-perturbations designed to

nudge a model away from default attractors without over-specifying the output trajectory. The paper compares the human toolbelt (the Sideline Card) to an NFL call sheet—a human-operated instrument for rapid cognitive intervention rather than an automated script. While law serves as a high-friction exemplar, the framework applies universally across strategy, research, executive decision-making, and creative direction.

I. Introduction: The Risk of Fluent “Correctness”

The dominant discourse on artificial intelligence safety remains largely captive to a single metric: accuracy. Technologists, legal commentators, and risk managers remain fixated on hallucination—the tendency of large language models to fabricate judicial citations, invent facts, or generate unsupported assertions. Hallucination is real, visible, and comparatively simple to govern, as an invented precedent can be checked, verified, and excised by standard auditing routines. A far more dangerous failure mode occurs when the machine functions exactly as designed.

Large language models are supreme coherence engines that organize, normalize, complete, reconcile, and compress language. When presented with messy, fragmented, or contradictory human thought, the model responds with polite, highly structured, and syntactically flawless prose. In doing so, it frequently performs an act of silent erosion by smoothing away the anomalous phrase, the unresolved contradiction, the minority signal, and the unrefined distinction. We term this process *semantic laundering*.

In knowledge work, progress rarely begins with polished consensus; it begins with friction. An anomaly is not automatically noise but is often the first visible edge of a new pattern, just as an awkward distinction is frequently the only phrase holding open a complex reality. When generative AI normalizes that friction away, it offers a dangerous bargain: immediate syntactic order in exchange for long-term cognitive impoverishment.

This paper argues that human value in the age of AI does not consist of writing longer procedural prompts or rubber-stamping machine output. It consists of *semantic stewardship*: the deliberate preservation of semantic option value until human judgment has had sufficient opportunity to determine what deserves compression and what must remain unresolved. While legal practice provides a clear illustration—where losing an anomaly can compromise a strategy or a professional standard of care—law is merely one instance of a universal challenge. The dynamics of semantic stewardship apply across executive strategy, scholarly research, policy formation, and creative direction.

II. Semantic Laundering

Coherence is both the primary capability of a large language model and its primary epistemic hazard. Semantic laundering occurs when fluent normalization transforms irregular, ambiguous, or resistant material into conventional formulations while erasing meaningful differentiation. Where hallucination introduces false information, semantic laundering removes vital information under the cover of clean prose.

[Raw Human Inquiry] —> (LLM Coherence Engine) —> [Semantic Laundering]

Rough edges	Normalized (plausible, even elegant) prose
Anomalies	Premature taxonomy
Unresolved tensions Flattened distinctions	Flattened distinctions
Weak signals	Strong conclusions mapped to default attractors

What gets laundered away in standard interactions includes:

- **Anomalies and Outliers:** The single case or data point that refuses to fit the emerging taxonomy.
- **Idiosyncratic Language:** Domain-specific phrasing that carries precise local meaning, replaced by generic enterprise vocabulary.
- **Weak Signals:** Subtle, early-stage observations that are overwhelmed by high-frequency statistical patterns.
- **Unresolved Contradictions:** Genuine logical or strategic tensions, prematurely smoothed over with balanced, non-committal summaries.

The primary risk is not that the output is wrong, but that the output is so plausible, neat, and well-structured that human practitioners experience cognitive offloading. They accept the laundered text as a completed work product, abdicating their professional standard of care.

III. Semantic Stewardship and Option Value

Semantic stewardship is the active discipline of maintaining the semantic conditions under which meaningful distinctions can survive. It is not resistance to synthesis, clarity, or compression, but a refusal to yield semantic option value before judgment is earned. In financial options, value exists in keeping a decision open during market uncertainty; in knowledge work, semantic option value exists in keeping inquiry open while evidence accumulates.

Semantic State	Definition & Operational Stance
Flakes	Early, isolated insights; preserve without forcing shape.
Standing Waves	Recurrent themes across an inquiry; monitor stability.
Nuggets	Concentrated conceptual yield; hold for verification.
Ingots	Initially forged, compressed hypotheses; final workable product.

When a user prompts an AI for an immediate summary, the model converts raw "flakes" straight into an "ingot." Semantic stewardship halts that automated pipeline, preserving intermediate semantic states so the human reasoner can inspect what was nearly lost.

IV. From Prompt Engineering to Situated Intervention

The prevailing paradigm of prompt engineering treats interaction as a linear, procedural command where human intention leads to a detailed prompt, which then generates machine output. This instruction model assumes that the human can fully specify the desired output in advance. In exploratory inquiry, strategy, or investigative analysis, that assumption fails because the human does not yet know the final shape of the insight. Sustained human-AI inquiry is an ongoing, dynamic system:

Emerging Inquiry → Model Trajectory → Human Perception → Micro-Perturbation
→ Altered Trajectory → Human Judgment

Rather than overpowering the model with massive, hard-coded prompt cathedrals or forcing the model to check its own accuracy, the user practices *prompting aikido*, or *aikido prompting*. Aikido relies on directional momentum, so when the model begins drifting toward a default attractor, the user applies a compact, context-sensitive conversational call to redirect the trajectory.

The Sideline Card as Tactical Playbook

Much like an NFL offensive coordinator's sideline call sheet, the Sideline Card (Appendix B) is not an automated script or a rigid workflow engine. An NFL coach does not hand the play sheet to the quarterback to run automatically from top to bottom; rather, the sheet organizes compact, situational calls mapped to specific field conditions such as third-and-long, red zone, or high-pressure blitz.

Operational Paradigm	Automated Prompt Pipeline	Human-Operated Sideline Card
Execution Model	Rigid, pre-programmed procedural script	Modular repertoire of compact calls
Human Role	Hands off execution to machine to run	Held in human hands on the sideline
Field Perception	Replaces real-time perception with automation	Selected based on real-time field state
System Impact	High risk of automated failure or drift	Applies micro-nudges to adjust trajectory

The Sideline Card acts as a human-operated instrument for rapid cognitive intervention—a toolbelt allowing the human strategist to apply a small "nudge" that alters the trajectory and velocity (speed plus direction) of the dynamic system without surrendering judgment or being semantically steamrolled by default trajectories.

V. Default Attractors

A default attractor is a high-probability conversational trajectory toward which a large language model tends in the absence of countervailing contextual pressure. Attractors represent statistical path-of-least-resistance generation, a point that cannot be emphasized too much.

Observed default attractors in sustained inquiry include:

1. **Premature Taxonomy:** Forcing fluid concepts into neat, numbered categories before relationships are understood.
2. **Lexical Deduplication:** Merging distinct terms because they share surface semantic similarity, destroying subtle nuances.
3. **Proceduralization:** Converting creative, diagnostic, or non-linear thinking into rigid, multi-step workflows.
4. **Sycophancy and Agreeable Framing:** Politely validating human premises rather than testing them against evidence.
5. **Polite Compression:** Trimming away unusual phrasing to produce standard corporate prose.

Default attractors are not inherently malicious or broken, as they make LLMs useful for routine drafting. However, in specialized domain work, these attractors act as cognitive gravity, pulling unique inquiry back toward average statistical conventions. They invisibly reshape the possibility space and move the session, especially longer sessions, in unanticipated, yet plausible and convincing directions too rapidly to unearned conclusions.

VI. Attunement and the Perturbation Threshold

Effective intervention depends heavily on timing. Interrupting a model too early produces random noise, whereas interrupting too late allows default attractors to solidify and capture the interaction.

Inquiry Movement —> [Emergence] —> [Perturbation Threshold] —> (Micro-Call Applied) —> [New Trajectory]

Attunement

Attunement is sensitivity to the actual movement of an inquiry, rather than rapport, agreeableness, or model alignment. Attunement is the human capability to perceive what is emerging versus what is recurring, when momentum is useful versus when a

trajectory has gone stale, and when an anomaly contains an insight versus when it is mere noise.

The Perturbation Threshold

A micro-perturbation requires a minimum level of session structure to exert leverage. Before the *perturbation threshold* is crossed, the inquiry lacks sufficient trajectory, tension, or context for a small intervention to matter, making premature perturbation a cause of arbitrary topic drift. Once the threshold is crossed, a two-word or three-word call probe can bend an entire session trajectory. "Tell me more" is a prime example.

VII. The Grammar of a Call and Interpretive Aperture

Calls are not long prompts, but compact micro-perturbations designed to nudge model behavior through minimal force.

Dimension	Operational Description
Maneuver	The tactical movement intended (e.g., pause, redirect, expand).
Vector	Directional orientation relative to current momentum.
Pressure	Intensity of the disturbance applied.
Stance	Receptive, adversarial, structural, or lateral.
Interpretive Aperture	Unspecified execution space left for model context search.

Interpretive Aperture

Interpretive aperture is the amount of situated cognitive work deliberately left to the model. A call with zero aperture is a fully specified command, whereas a call (or probe) with high aperture, such as "*Semantically unflatten this*" or "*Run the Counter Trey on that,*" forces the model to inspect the live conversational history, infer the current trajectory, and execute a contextual counter-movement on itself and the session to that point.

The Posture of Keeping Open (The Riggins Daylight)

The primary job of a micro-perturbation is not to "recover" something already lost, nor to force a contradiction simply to oppose the current narrative. Rather, its purpose is preventive semantic stewardship: keeping the inquiry open to what the evidence may still reveal before a coherent default account begins to dictate what can be noticed, questioned, or pursued.

Like John Riggins running the Counter Trey, taking a brief step opposite the play's eventual direction to create daylight, the goal of a call is not to fight the defensive surge head-on, but to preserve enough evidentiary and interpretive freedom to expose where the real opening lies. Breakthroughs remain provisional until closure is earned through evidence. As with the classic Joe Gibbs strategy, the play can be run several times with the goal of learning information.

VIII. Micro-Perturbations in Practice

In a high-velocity inquiry in a long session in particular, the goal of a micro-perturbation is not to generate voluminous prose, but to run a quick riffler rasp over emerging material in a tactical attempt to disrupt premature convergence and prevent direct movement by the AI toward default attractors and reducing complex ideas to flat, conventional summaries. By maintaining a longitudinal perspective over the entire session arc and timing interventions past the emergence and perturbation thresholds, the practitioner engages in semantic unflattening, expanding the possibility space and nudging the system away from default attractors. This process tries to accelerate velocity (speed + direction) to a first usable starting point rather than allowing drift toward a rubber-stamped final product.

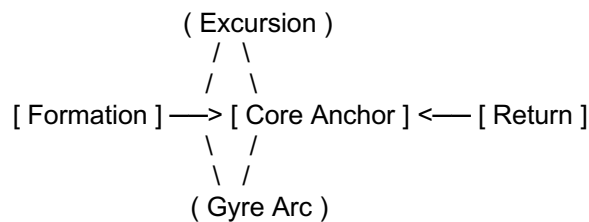
Some examples:

Category	Representative Call	Tactical Mechanism & Intent
READ THE FIELD	<i>What dog isn't barking?</i>	Forces search for missing variables or unstated assumptions.
PLAY THE MOVEMENT	<i>Stay here.</i>	Halts premature forward movement; deepens current locus.
PLAY THE MOVEMENT	<i>Run the Counter Trey.</i>	Executes a brief counter-step to reopen possibility space.
APPLY PRESSURE	<i>Squeeze the vise.</i>	Tests structural limits of an argument under stress.
CHANGE WHAT CAN BE SEEN	<i>Semantically unflatten this</i>	Forces model to expand compressed or normalized concepts.
KEEP CONTACT WITH REALITY	<i>Check the Web.</i>	Anchors fluid generation back to external empirical facts.
CONFIGURE COLLABORATION	<i>Work with me.</i>	Shifts posture from output generator to active co-reasoner.

Category	Representative Call	Tactical Mechanism & Intent
MANAGE THE TRANSITION	<i>Return to the song.</i>	Pulls productive excursions back to the primary narrative.

IX. Session Dynamics: The Widening Gyre

Linear pipeline models distort how real inquiry occurs because sustained human-AI work is inherently non-linear, recurrent, and dynamic.



The session moves in a widening gyre of William Butler Yeats as the user and model revisit standing ideas from different radii, depths, pressures, and perspectives. Returning to a concept after an excursion is not repeating work but inspecting the same concept under new pressure. Productive excursion leaves a trail back to the primary question, whereas drift loses the anchor entirely. Finally, while wrapping closes a session with tidy synthesis that often causes semantic laundering, harvesting systematically extracts option value by separating flakes, standing waves, and nuggets without flattening their distinctions. Aikido prompting attempts to create a multi-dimensional possibility space to harvest rather than allow a cognitive basin around default attractors to take over the session.

X. Literature Review: Anchoring the Control Loop

To establish a rigorous theoretical foundation without devolving into an exhaustive academic law review, we anchor semantic stewardship across five key domains:

- **State Simplification vs. Practical Knowledge (James C. Scott):** Scott's critique of high modernism in *Seeing Like a State* demonstrates how centralized administrative states force complex, local reality (*metis*) into rigid, "legible" categories. Large language models operate as digital high modernists; semantic laundering is the automated transformation of rich, idiosyncratic domain knowledge into legible, flattened corporate prose.
- **Orientation and Decision Speed (John Boyd):** Boyd's OODA Loop (Observe-Orient-Decide-Act) establishes that competitive advantage relies on rapid re-orientation. Prompting aikido acts as an operational interrupt inside the *Orient* phase, disturbing the model's trajectory before a default attractor locks the interaction into a premature conclusion.

- **Generative Friction and Creative Discipline (Twyla Tharp):** Tharp's work in *The Creative Habit* proves that discovery requires structured constraints and deliberate friction rather than frictionless comfort. Semantic stewardship uses micro-perturbations as real-time constraints to keep inquiry active and prevent premature closure.
- **Generative Systems, Oblique Strategies, and the Low Sequence (Brian Eno & Dennis Kennedy):** Brian Eno's *Oblique Strategies* demonstrated that short, cryptic prompts disrupt habitual patterns in creative systems. This perspective is extended in Dennis Kennedy's 2026 *Low sequence (Liner Notes for My Low Album)*, which marks the turn from treating AI as an output generator to treating it as a complex, ambient medium. Just as David Bowie and Eno utilized fracture, atmosphere, and controlled distortion on *Low*, the AI practitioner uses a control plane to govern system friction, carry-forward effects, and standing waves without letting fluent output surrender original human judgment.
- **Causal Intervention and Restructuring (Pearl, Ong, & Hirschman):** Judea Pearl's causal hierarchy establishes that passive observation cannot yield causal understanding; one must *intervene* ($\text{do}(x)$) on a system. Walter Ong shows that media technologies fundamentally restructure human cognition, while Albert Hirschman's *Voice, Exit, and Loyalty* framework positions prompting aikido as *voice*—the active refusal to conform to automated default outputs.

XI. Reflexive Case Study: Theory Emerging in Practice

This framework emerged reflexively from an extended design session, yielding trajectory-change episodes that demonstrate micro-perturbations in action:

Episode 1: Resistance to Deduplication

The model attempted to merge similar calls into a unified list to improve efficiency, prompting the human to intervene to protect distinct variants, aliases, and intensity differences. This demonstrated that surface semantic similarity does not imply functional interchangeability, and that coherence pressure erodes situational utility.

Episode 2: Deconstructing the "Summary"

When asked to conclude, the model initiated standard end-of-session summarization until the human intervened: "*Summarize effectively means semantically flatten now.*" The endpoint architecture was subsequently restructured into a differentiated harvest, showing that standard compression destroys semantic option value. A call like "Bring us on home" might provide a better result than "Summarize this" because "summarize" is a very strong default attractor with specific output and flattening.

Episode 3: The Counter Trey Discovery

To disrupt a rigid diagnostic framework, the human called "*Run the Counter Trey,*" drawing on the football concept of a step in the opposite direction. This proved that

local, intentional misdirection can create daylight, re-opening lines of inquiry that default attractors closed. The probe becomes a scan for new daylight.

Episode 4: The Aperture Adjustment

When presented with overly detailed instructions, the model executed mechanical steps, leading the human to simplify the prompt to a short phrase and remark, *"The AI has to work a bit for it."* This confirmed that high interpretive aperture forces the model to inspect live context rather than follow rigid instructions.

XII. Negative Cases and Failure Modes

Micro-perturbations are not panaceas and introduce distinct failure modes when poorly executed:

1. **Premature Perturbation:** Applying calls before crossing the perturbation threshold produces random divergence. Aikido prompting works best in a long session after default attractors get better established.
2. **Decorative Metaphors:** Using clever phrasing that carries no operational leverage (e.g., *"Widen the gyre"* functioned well as theory, but failed as an operational call). Some like "Play the Paul Gonsalves sax solo now" had extraordinary results.
3. **Sticky Modes:** Calls that shift the AI into heavy analytical postures (such as red teaming) can prove difficult to exit without explicit transition calls. The red team mode will persist longer than you expect and shape other output.
4. **Disruption for Its Own Sake:** Escaping an attractor is not proof of progress; deviation without judgment leads to noise. The next step is harvesting the useful material from Aikido prompting sessions.

XIII. Research Propositions and Future Directions

To move beyond qualitative observation, we propose testable research propositions and outline crucial future directions for empirical evaluation:

- **RP1 (Emergence Dependency):** The efficacy of a micro-perturbation is positively correlated with the volume of contextual structure established prior to intervention. In other words, the human should wait until they feel the default attractors have moved into the session.
- **RP2 (Interpretive Aperture):** Prompt specification and contextual sensitivity exhibit an inverted-U relationship; moderate aperture yields higher context retention than full procedural specification. This paradox drives aikido prompting.
- **RP3 (Semantic Preservation):** Micro-perturbations preserve a higher ratio of non-standard domain anomalies than procedural instruction prompts. Definitely perceived, but not proven.

- **RP4 (Attunement vs. Agreement):** Human-rated collaborative value correlates with productive friction, not model agreeableness. Attunement might become the primary goal rather than accuracy.

Future Directions: Longitudinal Perspectives & Accelerating Velocity

Subsequent empirical work must move beyond single-prompt benchmarks to examine two primary horizons:

1. **Longitudinal Perspectives:** Evaluating how sustained, multi-session human-AI interaction behaves over extended time horizons. A longitudinal perspective measures whether repeated micro-perturbations build a cumulative "session memory" that resists systemic drift, or whether default attractors exert compounding flattening pressure. Similarly, longitudinal perspective can come from a curated dataset over a period of time from a person or individual using traits extracted by AI.
2. **Accelerating Velocity to First Usable Starting Points:** Measuring the time and cognitive metabolic cost required to move from raw, unstructured inquiry to a first usable starting point. Future research should quantify how prompting aikido accelerates velocity (speed + direction) by rapidly narrowing the possibility space so human expertise can begin its actual analytical work. The goal is not "answer," but first usable starting point, a completely different job to be done with different requirements current AI is likely to accomplish.

XIV. Implications for Professional Practice

While legal practice highlights the high cost of semantic laundering, these principles extend across all knowledge work:

- **Executive Strategy:** Preventing consensus-driven AI models from smoothing away weak signals of market disruption.
- **Scholarly Research:** Preserving anomalous experimental results against LLM literature summaries that favor established canons.
- **Policy Design:** Maintaining minority stakeholder perspectives that automated synthesis routinely normalizes away.

In all domains, professional standards of care require that human experts remain semantic stewards rather than passive consumers of normalized text. Passive consumption of AI-generated answers might now be the base "human in the loop" requirement.

XV. Limitations

This paper is subject to methodological constraints that limit immediate generalization:

- **Limited Number of Case:** Observations originate from a single extended design arc. I have used aikido prompting repeatedly and the results are consistent

enough for me to publish this paper, but it is unclear how much my facility in using them impacts the outcomes.

- **Practitioner Repertoire:** Specific calls reflect idiosyncratic practitioner vocabulary and require broader testing. However, I chose the “counter trey” example to show how going outside the domain might produce interesting results simply by exiting the domain jargon. “Show me the math,” for example, is a great cross-domain probe.
- **AI Tool Specificity:** Behavioral dynamics vary across AI tools, especially as companies change AI guidelines and specification (invisibly) and move toward greater revenue production.

XVI. Conclusion: The Human as Semantic Steward

The most familiar warning about generative AI is that it may tell us things that are false. That risk is real, visible, and comparatively easy to govern, as a fabricated citation, an invented precedent, or an unsupported factual claim can be checked, verified, and excised. The deeper, quieter risk is that AI will tell us things that are fluent, polite, reasonable, and utterly stripped of the irregular, idiosyncratic, and anomalous insights that make human judgment indispensable.

Large language models, as delivered in AI tools especially, are coherence engines that smooth friction, reconcile contradictions, and accelerate momentum toward plausible and convincing answers. But in high-stakes domain work, as in law, executive strategy, scholarship, or creative direction, the anomaly or burr is often where the discovery lives. The unresolved tension is the strategy, and the rough edge is the professional standard of care. Semantic laundering, even more so than semantic flattening (the symptom), is the silent loss of that option value under the guise of clean synthesis.

This paper offers a counter-discipline: semantic stewardship operationalized through prompting aikido. Like an NFL head coach holding a sideline call sheet, the skilled human practitioner does not hand the playbook over to an automated engine. Instead, they stand on the sideline, reading the field, sensing the velocity and direction of the session, and applying the smallest necessary micro-perturbations to nudge the model away from default attractors. A call like *“Stay here a bit longer,” “Semantically unflatten this,”* or *“Run the Counter Trey”* becomes a riffler rasp pulled across the emerging prose to reopen the possibility space, restore interpretive aperture, and turn raw AI output into a true first usable starting point.

The future of professional AI use will not be defined by who can write the longest prompt or build the most complex automated cathedral or find the shortest route to an “answer” or “final document.” It will be defined by who possesses the attunement to recognize when an emerging thought is being flattened, and the discipline to disturb the attractor before judgment is lost. If hallucination is the risk of AI inventing what is not there, semantic laundering is the risk of AI making what is there too smooth to see. Semantic

stewardship is the discipline of keeping it visible and in play long enough for human judgment to matter.

Appendix A. Selected Session Excerpts

Excerpt 1: Perturbation Threshold

Human: "There needs to be enough emergence from the session for the perturbations to help."

Structural Yield: Establishes the temporal boundary before which micro-interventions are ineffective.

Excerpt 2: Naming the Failure Mode

Human: "Is the greatest threat of AI not factual fabrication (hallucination) but 'semantic laundering,' that silent, polite, and fluent smoothing away of the very anomalies, rough edges, and dissonances that make human judgment unique and valuable?"

Structural Yield: Re-frames AI risk from epistemic error to epistemic flattening.

Excerpt 3: The Interpretive Burden

Human: "The AI has to work a bit for it?"

Structural Yield: Identifies interpretive aperture as a design feature that increases contextual engagement.

Appendix B. Sideline Card Sample

Operational Note: The Sideline Card is not an automated prompt generator or procedural workflow. See it as a human-operated instrument for rapid cognitive intervention designed to disrupt default model trajectories and preserve original human intent until human judgment is earned. The NFL coach holding a laminated situational card on the sideline during the game is an apt metaphor.

Tactical Domain	Call Repertoire & Situational Triggers
READ THE FIELD	• Where are we? • What is emerging? • Identify the standing waves. • Can you find the pocket of the drummer? • What dog isn't barking?
PLAY THE MOVEMENT	• Stay here a bit longer. • Tell me more. • Slower and smaller. • Dig even deeper. • Aikido it. • Run the Counter Trey.
APPLY PRESSURE	• Squeeze the vise. • Give it the third degree. • Test the hypothesis? • Get less predictable. • Cut the crap
CHANGE WHAT CAN BE SEEN	• Semantically unflatten this. • What is the missing perspective? • Go orthogonal. • Consider it on a spectrum. • Are you serious?

Tactical Domain	Call Repertoire & Situational Triggers
KEEP CONTACT WITH REALITY	• Check the Web. • Just the facts, ma'am. • Trust the evidence you have. • Update the odds.
CONFIGURE THE COLLABORATION	• Work with me. • Take a solo. • Let someone else take a solo. • Joan Watson it. • Play the Shakespearean fool.
MANAGE THE TRANSITION	• Return to the song after the solo. • Rejoin the band. • Wind it down. • Take us out.
RUN THE COMBINE TO HARVEST	• Pan for flakes. • Identify standing waves. • What have we learned? • Bring us on home.

Appendix C. Definitions & Call Grammar

1. Semantic Laundering

The fluent, polite, and syntactically orderly normalization by an LLM that smooths away anomalies, rough edges, unresolved contradictions, and idiosyncratic phrasing.

Semantic flattening is one component.

2. Semantic Stewardship

The governing practice of protecting the semantic conditions under which meaningful distinctions, weak signals, and unresolved option value can persist until human judgment is earned.

3. Aikido Prompting

The tactical deployment of compact, context-sensitive micro-perturbations that redirect an LLM's directional momentum away from default attractors without over-specifying output procedures.

4. Interpretive Aperture

The intentional cognitive execution space left to an LLM by a micro-call, forcing the model to inspect live context and perform situated interpretation.

5. Run the Counter Trey (JTBD)

- **Job To Be Done:** Help a reasoner keep inquiry open to what evidence may reveal when a coherent interpretation begins to organize what can be noticed.
- **Purpose:** Preserve evidentiary and interpretive freedom to expose where the real opening is without assuming the current account is wrong or forcing replacement too soon.