

## MEMORANDUM FOR CONGRESSIONAL OVERSIGHT

**DATE:** July 13, 2026

**TO:** Professional Staff Director, Subcommittee on Oversight and Investigations, House Permanent Select Committee on Intelligence (HPSCI)

**FROM:** Ken Piggott, Retired Program Director (Cybersecurity/Advanced Technologies), U.S. Treasury ISSO

**SUBJECT:** Structural Mitigation of the "Confidence Problem" and Sensor-to-Policy Framing Failures in Strategic Intelligence Triage

### I. EXECUTIVE SUMMARY

As seen in recent global conflicts, objective analysis of intelligence data is an overwhelming task. It is reasonable to assume that Artificial Intelligence (AI) will be used increasingly for this analysis. The critical difficulty is that the reasoning and decision-making architecture remains buried deep within the LLM's non-transparent weights and code, resulting in an unvetted "Confidence Problem" at the policy handoff. To improve operational quality and transparency, the Recursive Critical Dialogue Protocol (RCDP) was developed. RCDP is a defined, open-source research model that recursively applies Systematic Scrutiny, Adversarial Steelmanning, and Explicit Falsifiability Constraints—evaluating the data cycle continuously until a new, actionable output is reached or a strict exit condition is met. This structured framework is designed to integrate into AI-assisted analytic pipelines to increase the rigor, epistemic vigilance, and governance of the process. RCDP is an open-source protocol available for download at [RCDP-AI.com](https://RCDP-AI.com), which also features The Kernel, its operational implementation designed to integrate these cognitive safeguards into structured human-AI workflows.

### II. WHEN CORRECT INTELLIGENCE STILL FAILS: TWO FAILURE MODES (Among Many)

The two failure modes below are illustrative, not exhaustive — drawn from a much longer list of ways institutions mishandle intelligence they already have. Both share a common shape: the underlying intelligence is not the point of failure. In the first, real data gets bent to fit what leadership already believes before it is ever reported up. In the second, the data is accurate and delivered on time, but the system built to act on it does not move. Either way, correct information still produces the wrong outcome.

**1. Cultural Echo Chambers & Data Falsification (The Bureaucratic Trap):** As observed in Russia's 2022 invasion of Ukraine, failures occur when intelligence is actively massaged to fit pre-existing executive narratives. Independent postwar analysis found that corrupt reporting practices had inflated official readiness figures while eroding real combat capability, and that senior military leadership is assessed to have concealed known readiness shortfalls from political leadership ahead of the invasion [1][2]. Mid-level reporting variations of this kind cause paper inventories to deviate catastrophically from actual battlefield readiness. Raw tracking data is rendered useless because analysts fear delivering dissenting facts to a rigid leadership structure.

**2. The Warning Paradox & Fluid Prose (The Policy-Intelligence Gap):** Conversely, advanced operations suffer from the Warning Paradox: the recurring pattern in which intelligence correctly anticipates a strategic event, but institutions fail to convert the warning into preventive action before it becomes a crisis [4]. Even when technical collection—such as advanced Measurement and Signature Intelligence (MASINT), localized thermal scanning, power-grid fluctuations, and RF emissions—is flawless, the system breaks down at the policy layer. Highly accurate strategic warnings regarding asymmetric resilience are often downplayed or bypassed by political decision-makers. Fluid, highly confident briefing prose routinely passes through the triage chain without a mechanism to track actual substantiation — a specific failure mode intelligence tradecraft has recognized since at least 1964, when imprecise qualitative language in estimates was first identified as a driver of misread warnings [5].

Both failure modes will only intensify as AI-driven analysis becomes more deeply embedded in the intelligence cycle. Automated systems can generate confident-sounding assessments at a volume and speed no human reviewer can independently verify in real time, meaning cultural echo chambers and the Warning Paradox will compound faster than they do today under human-paced production. This acceleration is precisely why claim-level governance

— enforced as each judgment is introduced, rather than reviewed once at the end of a product — becomes necessary rather than optional.

### III. PROPOSED SOLUTION: THE RCDP FRAMEWORK AS A COGNITIVE GOVERNOR

To insulate the Intelligence Community from these vulnerabilities, we must enforce structural friction between analysts and decision-makers. Moving boxes on an organizational chart or forcing wide-scale staff reductions will not fix the underlying cognitive flaws; it will only accelerate the creation of compliant echo chambers.

Instead, the briefing cycle must require three non-negotiable, protocol-driven guardrails before any strategic assessment shapes policy direction:

- **Mandatory Adversarial Steelmanning:** Analysts must explicitly formulate and defend the strongest possible counterargument to their own thesis at the exact moment a claim is introduced, rather than pooling alternative views as a minor appendix at the end of a report.
- **The Falsifiability Constraint:** No intelligence assessment may advance to an executive desk without an explicit failure condition. The briefing paper must clearly state: "Here is the specific, locatable piece of incoming data that would prove this assessment completely wrong."
- **Dynamic Divergence Checks:** At every major operational transition or mid-campaign shift, leaders must execute a formal framing review to answer: Are we adapting our methodology to the active threat, or are we running an advanced collection machine merely optimized to confirm our initial assumptions?

This proposal builds on existing practice rather than substituting for it. The Intelligence Community already mandates a form of adversarial review under Intelligence Community Directive 203 (ICD 203), Analytic Standards, Tradecraft Standard 4, "Incorporates analysis of alternatives" [3]. RCDP's contribution is not to introduce oversight where none exists, but to move that standard from a pooled, end-of-product exercise to a claim-level, timestamped requirement enforced at the moment a judgment is introduced. This is offered as a testable proposition, not a settled result: it should be piloted against declassified case material and evaluated on whether claim-level steelmanning surfaces objections that end-of-product alternative analysis misses. If it does not outperform existing tradecraft practice under that test, the proposal fails on its own terms.

### IV. ACADEMIC VALIDATION AND METRICS DISCLOSURE

The RCDP methodology is not merely a theoretical framework; it is an active, structured research instrument currently in final review at Frontiers. The protocol explicitly answers a graduated four-question AI-disclosure framework proposed within that same body of research — not an external benchmark — covering task execution, structural thresholds, epistemic authority, and context document reproducibility. Peer review is intended to allow the proposed safeguards to be independently audited and stress-tested before any claim is made about performance at scale; this memorandum should not be read as asserting that the methodology has been validated for intelligence-community use, only that it is available now for evaluation.

### V. SUPPORTING DOCUMENTATION AND REFERENCE

The foundational theoretical mechanics, version governance, and operational parameters of the Recursive Critical Dialogue Protocol (RCDP) are open-source and managed as a live human-AI collaborative research initiative. RCDP is also packaged as an installable AI "skill," allowing direct incorporation into existing agent-based analytic pipelines without requiring custom integration work. Full documentation, framework downloads, and active methodology findings can be verified directly at the project's central hub: <https://rcdp-ai.com>

### VI. A STATEMENT OF GOOD FAITH

The most direct way to evaluate RCDP is to test it directly: copy the mini-Kernel document from the project hub into the start of any AI session — professional or personal — and observe the difference in analytic behavior firsthand. There is no proprietary mechanism being withheld; the full protocol is open and reproducible.

This material is submitted in the same spirit as this author's prior public service: a Treasury Department appointment as an Information System Security Officer, taken under oath to the Constitution, adjudication for a Public Trust position, and membership in InfraGard, which requires each member to pledge to abide by a defined code of ethics. These commitments are noted here not as credentials, but as context for the good faith underlying this submission.

## REFERENCES

- [1] Renz, B. (2023). Western Estimates of Russian Military Capabilities and the Invasion of Ukraine. Problems of Post-Communism. <https://doi.org/10.1080/10758216.2023.2253359>
- [2] Radio Free Europe/Radio Liberty. A General's Correspondence Details the Wartime Role of Graft in the Russian Military. <https://www.rferl.org/a/russia-military-corruption-ukraine-war-bribery/33692902.html>
- [3] Office of the Director of National Intelligence (2015). Intelligence Community Directive 203: Analytic Standards. [https://www.intelligence.gov/assets/documents/intelligence-community-directives/ICD\\_203.pdf](https://www.intelligence.gov/assets/documents/intelligence-community-directives/ICD_203.pdf)
- [4] Robert, E. J. (2026). The Warning Paradox: Why Correct Intelligence Often Fails. The Cipher Brief. <https://www.thecipherbrief.com/why-correct-intelligence-often-fails>
- [5] Kent, S. (1964). Words of Estimative Probability. CIA Center for the Study of Intelligence. <https://www.cia.gov/resources/csi/static/Words-of-Estimative-Probability.pdf>

Respectfully submitted,

**Ken Piggott**

*Retired Program Director, Cybersecurity and Advanced Technologies, U.S. Treasury ISSO*

*Canton, Michigan*

*ken@krp.biz | (313) 460-1540*