

Human × AI Collaboration — Running the Gate With a Machine (6 sessions)

The doctrine you are learning is also executable by a machine — and you still own the decisions. This module is the taught realization of literacy/how-the-machine-learns.md: it turns the concept layer into a working skill. By the end you will not just know the machine exists; you will be able to brief it, read its output critically, hold it accountable, and run the Dream Gate as the accountable member of a two-operator team.

Audience: everyone, after the shared core. **Prerequisites:** shared-core (or at minimum doctrine/01–08) and a first read of literacy/how-the-machine-learns.md.

Readings: doctrine/03, doctrine/04, doctrine/05, doctrine/08, doctrine/09, literacy/how-the-machine-learns.md, and the drill bank at practice/human-ai-drills.md (assigned with the sessions).

Session 1 — The two-operator team

The machine can run the pipeline. It cannot own it.

Learning objectives. By the end of this session, students can: (1) explain what a human professional and an AI agent each do best in capture and proposal; (2) state the division-of-labor rule — the human is accountable, the agent is the instrument; (3) classify any pursuit task by the operator that should do it.

Session plan (60 min). - Open (5 min): “What can a machine do better than you in this profession? What can you do better than a machine?” — two lists on the board. - Teach (20 min): the two-operator team. What the machine is reliably better at; what the human is irreplaceable at; why “the machine proposes, you dispose” is the spine of doctrine/08. - Apply (20 min): the operator map. Give teams a list of real pursuit tasks; classify each M (machine best), H (human best), or H+M (co-produced) with one line of reasoning. - Discuss (10 min): “Where does an organization go wrong when it treats the machine as a colleague instead of an instrument?” - Close (5 min): assignment and preview of the module.

Body. Your profession has always been a team sport. The new teammate is not a colleague with a vote — it is an **instrument**: an AI agent that reads every page, never sleeps, and produces professional-looking output in minutes. It is like a brilliant analyst who has no skin in the game, no reputation to protect, and no judgment beyond what its knowledge library (the corpus) and your briefing give it. The discipline of working with it is the discipline of knowing, for every task, which operator should do it.

What the machine is reliably better at. The machine wins on the deterministic, the verifiable, and the well-scoped: sweeping a 90-page solicitation for every

“shall,” checking page limits and format rules, computing the score arithmetic, retrieving a known fact from its corpus, harvesting agency win patterns at scale, and drafting a credible first pass against the evaluation weights. These are the jobs where a tired human at 2 a.m. is the weakest link. The machine is the best compliance officer and first-draft author you will ever staff — because it is exhaustive where you are fallible and fast where you are slow.

What the human is irreplaceable at. The human owns the judgment, the novelty, and the irreversible. A machine cannot decide whether this pursuit fits the portfolio you are actually building. It cannot read the room — the unstated preference in a customer briefing, the signal that the incumbent is quietly in trouble. It cannot carry the credibility that a lived relationship puts behind a narrative. It cannot touch the integrity boundary — the organizational-conflict-of-interest judgment, the data-rights call, the ethics line (doctrine/01, doctrine/02). And it must never be the one who signs.

The division-of-labor rule. This is the same rule doctrine/08 teaches about tools, applied to the machine: the machine is the tool; the doctrine is the shared language; the professional is the owner. Concretely: **the machine proposes; you dispose.** It can draft, score, and warn. It does not decide, and it does not sign. Every gate that commits real resources — bid/no-bid, threshold, readiness, submission — keeps a human owner who is a person (doctrine/04). That is what the Dream Gate means in the machine age: the machine can fill the whole pipeline with proposals, and the human still runs every gate.

Accountability. The two-operator team has one captain. If the machine drafts a section with a fabricated reference and it goes to the agency, the error is yours — because the machine is an instrument and the instrument answers to the operator. “The machine wrote it” is not a defense in this discipline; it is an admission that you stopped being a professional. The day you stop checking the machine’s work is the day you stop being the accountable member of the team.

Worked example. A NOFO lands on a Friday at 4 p.m. Two operators, one pursuit. The machine, overnight: pulls the solicitation, checks eligibility against the record (NAICS alignment, size standard, set-aside, active registration, the research-partner work-share), extracts the three deadlines, builds a compliance skeleton of fourteen “shall”s, and leaves a first-pass technical outline plus a corpus note on what has won at this agency family. The human, Monday morning: a fifteen-minute strategy read — does this fit the portfolio? Can the firm staff it as its one active pursuit? Is it actually positioned to win, or is the machine’s compliance sweep just making a bad idea look organized? The human runs the qualify gate and makes the call. The machine did a week of analyst work in a night; the human made the decision that justified the week. Neither could have done the other’s half.

Discussion prompts. 1. “The machine can run a pursuit; it cannot own it.” Defend or attack. 2. Name one task where delegating to the machine saves real money, and one where it would be malpractice. 3. When does a machine

proposal stop being a draft and start being a decision — and who moves it from one to the other?

Homework / reading. Read doctrine/08 and literacy/how-the-machine-learns.md. Deliverable: a one-page **operator map** for one pipeline stage of your choice — list its tasks, mark each M / H / H+M, and give one line of reasoning per row.

Self-check. 1. What three kinds of work is the machine reliably better at? 2. What does “the machine proposes; you dispose” mean at a gate? 3. Why is “the machine wrote it” not a defense?

Session 2 — Briefing the machine

The machine cannot capture what the room knows but the brief does not.

Learning objectives. By the end of this session, students can: (1) explain the garbage-in principle — a vague or wrong brief produces confident, unusable output; (2) write a structured capture brief with all five parts; (3) diagnose the failure mode of an under-specified brief.

Session plan (60 min). - Open (5 min): “What would you tell a brilliant new analyst who knows nothing about this firm, before they draft a section?” — that list is the brief. - Teach (20 min): the five parts of a capture brief; the garbage-in failure modes (vague, confident-wrong, silent). - Apply (25 min): in pairs, write two briefs for one real, closed solicitation — a vague one and a structured five-part one — and compare what each would produce. - Discuss (5 min): “Which part is most often missing in a real brief, and what does its absence produce?” - Close (5 min): assignment.

Body. The machine is a mirror of two things: its corpus and your briefing. You control the second. The **capture brief** is the written rendering of the capture plan handed to the machine — the same handoff doctrine/09 teaches between capture and proposal, now between human and machine. A good brief makes the machine’s output usable; a bad brief makes it confident and wrong. This is the discipline of **garbage in, garbage out** — with a machine twist: the machine will not hand back garbage politely, it will hand back garbage *fluently*, formatted like a professional, and full of invented specifics that paper over the gaps in what you gave it.

The five parts of a capture brief.

1. **The opportunity record.** The solicitation itself — not your paraphrase. The NAICS code, the set-aside, the agency, the value, the three deadlines. The machine reads the document; give it the document.
2. **The constraints.** The rules of the road: page limits, format, the evaluation weights, the threshold, the required certifications, and what is off-limits (“do not draft the past-performance volume yet,” “do not mention the

teaming partner”). Constraints are as important as content — they are what stop the machine from proposing a 60-page answer to a 15-page limit.

3. **The evidence.** The firm’s real delivered work, its certifications, its partner, its record — drawn from the record, not invented. The machine will use what you give it; if you give it nothing, it will borrow from the corpus, and borrowed past performance is fabricated past performance.
4. **The stance.** The win strategy if one exists: the named themes, the candidate discriminators, the price posture, the ghost-theme targets written as observations, not names. If you have not decided the stance yet, say so — a brief that pretends to have a strategy when it does not produces a proposal that pretends too.
5. **The questions.** What you want back and in what form: a compliance matrix, a first-draft section, a pWin factor sheet with an evidence column, a gap list. The machine is a tool; tools take instructions, not hopes.

The garbage-in failure modes. Three, and a professional can name all three. The **vague brief** (“help us bid on this”) returns boilerplate — generic themes any firm could claim, a compliance matrix that misses the exotic “shall”s. The **confident-wrong brief** feeds the machine a wrong deadline, a wrong NAICS, a wrong set-aside — and the machine propagates the error at full confidence, because it has no reason to doubt the operator. The **silent brief** leaves out the one thing the room knows: the customer signal that the incumbent is in trouble, the no-new-hires constraint, the true cost floor. The machine will not ask; it will fill the silence with the corpus. A good brief is a forcing function for clarity: if you cannot brief a machine, you cannot brief a colleague — because the brief is just the capture plan written down.

Worked example. Two briefs for the same STTR-shaped NOFO (the one from Session 1). The **vague brief**: “We should bid on this energy digital-twin NOFO. Draft the technical approach.” The machine returns generic win themes, a compliance matrix that misses the 40/30 work-share “shall,” and a technical approach any firm could have written. The **structured brief** gives the machine the opportunity record (the NOFO, the topic, the \$250,000 fixed award), the constraints (15-page technical volume, the 40/30 work-share, the PI-primary-employment rule, the 0.42 threshold), the evidence (the bench of physicists and AI specialists, the four NAICS codes, the INSTAR Lab partnership), the stance (candidate discriminator: physics-model pedigree plus a research partner; ghost target: larger primes with no research institution), and the deliverables (a compliance matrix, a first-pass technical outline that mirrors the evaluation weights, and a pWin factor sheet with an evidence column). The machine returns a usable first pass: all eleven “shall”s mapped, an outline that spends its pages where the weights are, and a factor sheet that flags past performance as the honest weak spot. The difference in the output was entirely the difference in the brief.

Discussion prompts. 1. Which of the five parts is most often missing in a real brief, and what does its absence produce? 2. Defend: “if you cannot

brief a machine, you cannot brief a colleague.” 3. Why is “give the machine the document, not your paraphrase” a rule and not a suggestion?

Homework / reading. Read doctrine/09 (the capture handoff) and the compliance-reading section of doctrine/02. Deliverable: write a one-page capture brief for a real, closed solicitation from SAM.gov or Grants.gov — all five parts — written as if briefing a machine operator who knows only the corpus.

Self-check. 1. Name the five parts of a capture brief. 2. What is the confident-wrong failure mode, and how does the machine make it worse? 3. Why must you give the machine the document, not your paraphrase?

Session 3 — Reviewing machine output

A fluent, confident machine output is not a correct one.

Learning objectives. By the end of this session, students can: (1) apply the D8.5 discipline — *is it right, not just plausible*; (2) run the five-point review checklist on any machine output; (3) recognize the three hallucination signatures and verify claims against the source solicitation.

Session plan (75 min). - Open (5 min): “What is the difference between a machine output that *looks* professional and one that *is* correct?” - Teach (25 min): D8.5 and the five-point review checklist; the three hallucination signatures. - Apply (35 min): the review clinic. Hand teams a machine-drafted section with planted errors (draw on the drill bank) and run the checklist; findings written as *claim + evidence + verdict + required change*. - Discuss (5 min): “Which checklist item catches the most expensive errors, and why?” - Close (5 min): assignment.

Body. The machine writes like a professional and can be wrong like a confident amateur. The core skill of the two-operator team is therefore **review** — the literacy page calls it D8.5: *is it right, not just plausible*. A fluent, confident machine output is not a correct one. The machine does not bluff in the human sense; it *completes the pattern*. Give it a sentence that needs a citation, and it supplies one that looks right. Give it a claim that needs evidence, and it manufactures the shape of evidence. That is exactly why fluency is the risk: a sloppy error is easy to catch, but a confident, well-formatted invention is built to pass.

The five-point review checklist. Run it on every machine output you intend to act on.

1. **Source check.** Every factual claim is traceable to the solicitation or the corpus. If the output cites a section, a page, or a clause, open that source and read it.

2. **Compliance check.** Does it cover every “shall”? Run the machine’s compliance matrix against your own read of the document — not against the machine’s summary of the document.
3. **Evidence check.** Is every claim backed by evidence? A win theme is a promise; the evidence is the proof. A discriminator needs proof the competitor cannot make the same claim.
4. **Discriminator check.** Is the “discriminator” actually unique and defensible, or a baseline dressed up in a suit? Does it move the score, or just establish that you are qualified?
5. **Reference check.** Do the cited references exist — the FAR clause, the page number, the quote, the contract number, the CPARS rating, the customer name and phone?

The three hallucination signatures. Learn to smell them. **Fabricated references** — a FAR clause that does not exist, a quote that is not in the solicitation, a contract number that was never awarded. **Unsupported claims** — a “discriminator” with no evidence, a win theme with no customer basis, a capability assertion with nothing behind it. **Invented facts** — past performance that never happened, a certification the firm does not hold, a deadline that is wrong. Each one passes the “reads well” test and fails the “is it true” test.

Ground truth. When machine output and the source solicitation disagree, the document wins — every time. “The machine said so” carries no weight against the published text. This is the same discipline doctrine/02 teaches for reading RFPs: the compliance reading answers *what is required*; the review of machine output answers *is this requirement real, and is it satisfied*. The source solicitation is the one ground-truth document in the room.

Worked example. A machine-drafted technical section for the Session 1 NOFO. Running the checklist finds: (1) the section cites **FAR 52.226-13** for a data-integration set-aside — a clause that does not exist; **source check fails**. (2) It asserts “Ravonics has delivered digital-twin models to federal customers” — the record we gave the machine shows laboratory-affiliated work, not federal digital-twin delivery; **evidence check fails**. (3) It quotes Section M as requiring “proven technology preferred” — the NOFO actually says it is funding *proof-of-concept* research and explicitly welcomes unproven ideas; **reference check fails**, and the misquote would have pushed the proposal in exactly the wrong direction. Three findings, three overrides, each written as *claim + evidence + verdict + required change*. The document won all three.

Discussion prompts. 1. Why is fluency the risk — what makes a confident wrong answer harder to catch than a sloppy one? 2. Which checklist item catches the most expensive errors, and why? 3. What would happen to a firm that reviewed machine output for format but never for truth?

Homework / reading. Read doctrine/02 and doctrine/05. Deliverable: take any machine output you can find (or drills HA-01 through HA-03 from practice/human-ai-drills.md), run the five-point checklist, and produce a review

memo with findings written as *claim + evidence + verdict + required change*.

Self-check. 1. What are the three hallucination signatures? 2. When machine output and the solicitation disagree, who wins? 3. What does “a baseline dressed up as a discriminator” look like in a machine output?

Session 4 — Holding the machine accountable

The machine proposes; you dispose. And you write it down.

Learning objectives. By the end of this session, students can: (1) explain the override log and what it is for; (2) audit an agent’s rationale — demand sources, chain, and confidence; (3) run a red-team pass (black-hat and ghost) on machine output.

Session plan (60 min). - Open (5 min): “How do you hold an instrument accountable? What would proof of accountability look like?” - Teach (20 min): the three tools — the override log, the Dream Gate, and the red team. Auditing a rationale: claim → source → conclusion → confidence. - Apply (25 min): the gate-review exercise. Give teams a machine score recommendation with a fabricated basis (draw on drill HA-06); they audit the rationale, re-score, and log overrides; then a black-hat pass on a machine theme. - Discuss (5 min): “What would an override log look like if nobody ever overrode the machine — is that good or bad?” - Close (5 min): assignment.

Body. Accountability is a practice, not a principle — and in a two-operator team the practice has three tools: the **override log**, the **Dream Gate**, and the **red team**.

The override log. Every time you change, reject, or grudgingly accept a machine output, you record it: the machine’s claim, the evidence you found, your action, your rationale. Four columns — nothing more. The log is not bureaucracy; it is the **calibration record**. Overrides reviewed over time tell you where the machine is reliably wrong — and where *you* are reliably wrong, because the machine catches your drift and your blind spots as often as you catch its hallucinations. The log feeds the learning loop (pipeline stage 9): the machine’s pWin estimates get checked against real outcomes, its calibration gets corrected, and the human’s override patterns become the machine’s next lesson. A team that overrides and does not log it is a team that does not learn.

The Dream Gate. The machine can draft, score, and warn. It does not sign. Every gate that commits real resources — bid/no-bid, threshold, readiness, submission — keeps a human owner. That is doctrine/04’s gate discipline applied to the machine age: criteria written in advance, the owner makes the call, and the owner is a person. The gate is where accountability bites — not because the machine is untrustworthy, but because the decision is yours either way. The gate is the difference between “the machine recommended it” and “I decided it.”

Auditing the rationale. When you are about to act on a machine call, make it show its work. What source did you use? What did the source say? What did you assume? What is your confidence, and why? Demand the chain — claim → source → conclusion — and verify the source yourself. A rationale that cannot be checked is a guess. A pWin of 0.58 with a factor sheet that cites fabricated past performance is not a score; it is a decorated guess. The audit is the moment the machine’s fluency meets your scrutiny.

The red team. The critics your profession already staffs for proposals (doctrine/09) apply to machine output. The **black-hat pass** assumes the machine’s output is wrong and tries to break it: attack the discriminators, attack the price, attack the compliance, attack the assumptions. The **ghost pass** checks the “ghost themes”: do they actually target a competitor without naming it, and do they have capture intelligence behind them — or are they hollow posturing? A machine output that survives the black-hat and ghost passes is a machine output worth a gate.

Worked example. The gate meeting. The machine recommends **BID** on a \$250,000 pursuit: pWin 0.58, composite \$145,000, above the 0.42 / \$105,000 threshold. The capture lead audits the rationale. The machine’s pWin factor sheet weights past performance at 0.85 — because it pulled a “similar engagement” from the corpus and presented it as Ravonics’s own record. It is not in the record; it is a fabricated basis. The human re-scores past performance at 0.45. pWin drops to 0.44 — still above the 0.42 line, but the call is now **BID WITH CONDITIONS**, not BID. Override logged. The red team then runs a black-hat pass on the machine’s technical approach and finds the “ghost theme” names a competitor outright. Second override logged. The pursuit proceeds with conditions, and both overrides go to the next portfolio meeting as calibration data.

Discussion prompts. 1. What would an override log look like if nobody ever overrode the machine — is that a good sign or a bad sign? 2. Defend: “auditing a machine’s rationale is the new compliance.” 3. Why does the gate need a person when the machine can be trained to follow the criteria perfectly?

Homework / reading. Read doctrine/04 (gates) and doctrine/09 (color teams). Deliverable: create a one-page override-log template and fill it with three entries from any machine output you have reviewed this week — real or from the drill bank.

Self-check. 1. What are the four columns of the override log, and what is the log for? 2. Why is the Dream Gate where accountability bites? 3. What does the black-hat pass do, and why apply it to machine output?

Session 5 — The machine’s knowledge is not yours

The corpus tells you what has won. You tell it what will win.

Learning objectives. By the end of this session, students can: (1) distinguish the corpus — a record of the past — from lived judgment — a position on the future; (2) state when to trust the corpus and when to override it; (3) recognize bias and echo-chamber risk in machine output.

Session plan (60 min). - Open (5 min): “What is the difference between knowing what has won and knowing what will win?” - Teach (20 min): the corpus as a record, not a prophecy. When to trust it; when to override it. Bias and the echo chamber. - Apply (25 min): the trust/override exercise. Give teams three machine win themes (one corpus echo, one transferred pattern, one grounded in the actual pursuit) and have them sort, justify, and log overrides. - Discuss (5 min): “What does an echo-chamber machine output look like in a proposal?” - Close (5 min): assignment.

Body. The corpus is the machine’s knowledge library — past performance, winning themes, agency win patterns, the rules and playbooks of the machine rendering. It is a record of **what has worked**. It is not a guarantee of **what will work**. The single most important judgment you will make as the human operator is telling the two apart. The corpus is a consultant with a perfect memory of yesterday; you are the professional who must answer for tomorrow.

When to trust the corpus. Trust it on the deterministic, the verifiable, and the well-scoped: the eligibility rules and compliance floors (encoded public rules, checkable), the arithmetic, and the retrieval of a known fact from the record. Trust it on **winners-by-agency patterns** — but only as *hypotheses* to test against this pursuit, never as verdicts. Trust it on past-performance narratives — after you have verified that the delivered work is actually in the record. The corpus is strongest exactly where the machine is strongest: well-scoped and grounded.

When to override. Override on **novelty** — a competitive situation the corpus has never seen; the corpus has no answer and will confidently improvise one. Override on **stale patterns** — what won three years ago in a different pool is data, not destiny. Override on **transferred patterns** — the theme that worked for a large prime does not transfer to a five-person HUBZone firm, and the machine will transfer it anyway because it reads patterns, not identity. And override any claim the corpus presents as fact that is actually a correlation: “this theme appeared in winning proposals at this agency” is not “this theme is why they won.” The corpus records what coincided; you must judge what caused.

Bias and the echo chamber. The corpus is built from the past, so it reproduces the past — including the past’s assumptions. If it has read mostly winning narratives from large primes, it will echo incumbency assumptions and generic “deep bench” themes even in a set-aside for small firms. That is the echo-chamber risk: the machine does not know it is inside a pattern; it only knows the pattern. The human is the one who breaks the echo — with the customer relationship, the room’s knowledge, the novel signal, the unstated preference the corpus cannot read. Never let a pattern from the corpus override a fact about this pursuit.

The one-line rule. The corpus tells you what has won; you tell it what will win. When the machine proposes a theme because it is *familiar*, and you override because it is *wrong for this pursuit*, that is not a disagreement — that is the two-operator team working as designed.

Worked example. The machine proposes a win theme for Ravonics on a HUBZone set-aside: “a deep bench of senior engineers with decades of combined federal experience” — a corpus agency-win pattern that worked for a large incumbent. Ravonics is five people, with a thin federal award record, a HUBZone certification, and a research partner. The theme does not transfer — and worse, it is a baseline any large firm could claim. The human overrides with a theme grounded in the actual pursuit: “a HUBZone small business paired with a research institution, delivering publishable proof-of-concept physics — the exact shape this NOFO funds.” The override log entry records both: the machine’s theme rejected (baseline + untransferable pattern), the human’s theme adopted (grounded in the NOFO and Ravonics’s real edge). The corpus was not wrong — it was in the wrong room.

Discussion prompts. 1. When is a pattern from the corpus a fact, and when is it a correlation dressed as a fact? 2. What does an echo-chamber machine output look like in a proposal, and how do you catch it? 3. Why is “the corpus has always done it this way” the exact sentence a professional should distrust?

Homework / reading. Read doctrine/07 (the lifecycle — why past performance is evidence, not destiny) and re-read the D8.5 section of literacy/how-the-machine-learns.md. Deliverable: a one-page “trust the corpus, override the corpus” memo for one pursuit — three corpus claims you would trust and three you would override, one line of reasoning each.

Self-check. 1. Why is the corpus a record of the past, not a plan for the future? 2. Give one situation where a corpus win pattern should be overridden. 3. What is the echo-chamber risk, and who is responsible for breaking it?

Session 6 — The team at work

The golden thread runs through both of you — the same doctrine, practiced on the same case.

Learning objectives. By the end of this session, students can: (1) walk a full pursuit end-to-end with a human and a machine, naming who does what at each stage; (2) run a color-team pass with a human running the gate; (3) describe what excellent human-agent teamwork looks like as a repeatable pattern.

Session plan (75 min). - Open (5 min): “Where is the machine strongest in a pursuit — and where does the human’s judgment change the outcome?” - Teach (20 min): the co-production pattern — the operator table across all nine pipeline stages; the gate rule; the override discipline; the golden thread. - Apply

(40 min): the full scenario. Give teams the compact NOFO; run it hour by hour — qualify, score, capture plan, color-team pass, gate — with one player as the machine operator and one as the accountable human. - Discuss (5 min): “What would the same pursuit look like with no human in the loop, and where would it fail?” - Close (5 min): bridge to the capstone and the drill bank.

Body. Pull it together. A pursuit with a human and a machine is not the human’s pursuit with a helper, nor the machine’s pursuit with an overseer — it is a **co-production** where each operator works where it is strongest, and a human runs every gate. Here is the pattern, across the pipeline (doctrine/03):

Stage	The machine does	The human does	The gate
Sense	Surfaces and pre-screens the opportunity; extracts the identifiers and deadlines	Decides it is worth a closer look	qualify
Qualify	Runs eligibility: NAICS alignment, size standard, set-aside, active registration, work-share	Makes the qualify call — fit, portfolio, capacity	qualify
Score	Assembles the pWin factor sheet from the brief and the corpus, with an evidence column	Argues the factors, sets the value, runs the sensitivity, owns the threshold call	threshold
Pursue	Drafts the ORBITAL skeleton from the score and the brief	Fills the judgment axes, reviews all seven, decides to commit	threshold
Design	Drafts first passes of the volumes to the evaluation weights	Sets the win strategy and themes, edits, writes the credibility-bearing narrative	review

Stage	The machine does	The human does	The gate
Audit	Runs the compliance sweep; produces the color-team materials	Runs the color teams (pink, red, gold, black-hat on the machine's own output); owns the override log	review / readiness
Submit	Assembles the package; re-checks the mechanical facts (page count, format, portal)	Verifies; authorizes; a person submits or signs	readiness
Deliver	Tracks milestones, flags deviations	Runs the delivery; the delivered work becomes the next pursuit's past performance	—
Learn	Captures the outcome, updates calibration, logs the post-mortem data	Leads the post-mortem, updates the pWin factors, reviews the override log	—

The pattern of excellent teamwork. Four marks. First, **each operator works where it is best** — the machine never makes a judgment call it is not equipped for, and the human never hand-checks a compliance sweep item by item. Second, **every gate has a human owner** — the machine can assemble the evidence, but the criteria are applied by a person (doctrine/04). Third, **every override is logged** — the calibration record is the team's memory. Fourth, **both operators practice the same doctrine on the same case** — the golden thread. When the human capture team and the machine pursuit layer both run Ravonics through the same pipeline, “gate,” “pWin,” and “ORBITAL” mean the same thing to both — that shared meaning is what makes the handoff clean instead of a translation.

The warning signs. A machine that is never overridden is a machine that is never audited. A human who never logs an override is a human who has stopped thinking. The failure state of the two-operator team is not the machine running away — it is the human rubber-stamping. The gate exists to make that boring: criteria written in advance, owner makes the call, the call is written down.

Worked example. The full scenario, in the shape of the worked exemplar. A compact STTR-shaped NOFO arrives on a Friday. **Morning — qualify.**

The machine pre-screens overnight: eligible (NAICS fits, HUBZone applies, the research-partner work-share is satisfiable), all three deadlines extracted. The human reads the fifteen-minute strategy pass and qualifies — this fits the one-pursuit portfolio. **Midday — score.** The machine assembles the factor sheet; the human argues the factors — past performance honestly weak, technical approach strong, topic fit near perfect — and runs the sensitivity. $pWin$ $0.62 \times \$250,000 = \$155,000$, comfortably over the $0.42 / \$105,000$ threshold. **Afternoon — capture plan and color team.** The machine drafts the ORBITAL skeleton and first-pass volumes to the weights; the human sets the win strategy and the two discriminators. The red team (fresh eyes, including a black-hat pass on the machine’s own draft) finds the machine’s ghost theme names a competitor — override logged, rewritten as a true ghost. The pink team finds the machine’s compliance matrix missed the PI-primary-employment “shall” — override logged, row added. **Evening — the gate.** The human runs the readiness gate against the criteria written in advance: position built, teaming solid, compliance complete, price supported. **BID.** The override log has three entries, the calibration file has three lessons, and the submission goes in with a person accountable for every word of it.

Discussion prompts. 1. Where is the machine strongest in this scenario, and where did the human’s judgment change the outcome? 2. What would the same pursuit look like with no human in the loop, and where would it fail? 3. Why does practicing on the same case — the golden thread — make a human and a machine better teammates?

Homework / reading. Re-read the worked exemplar (practice/worked-exemplar-day-in-the-life.md) and doctrine/03. Deliverable: run the co-production pattern on a real, closed solicitation — produce the operator table for its pipeline, the gate calls with the human owner named, and an override log with at least three entries.

Self-check. 1. Name one machine action and one human action at each of the nine pipeline stages. 2. What makes a gate human-owned? 3. What does a year of override logs reveal about a team?

Capstone integration

This module is the taught realization of literacy/how-the-machine-learns.md, and the practice layer exercises it. The drill bank at practice/human-ai-drills.md is where you review machine output containing planted errors — find each error, say why it is wrong, and log an override. The two AI-accountability challenges (CH-23 and CH-24 in practice/challenges) are the judgment layer: full scenarios that drop you into a working situation with a machine operator and force the same discipline under pressure. When you reach the capstone (course/capstone-build-an-orbital.md), build the ORBITAL as a two-operator team — brief the machine, review its drafts, run the color-team pass, and hand in the override log

as part of your defense. The machine can run the pipeline; the capstone proves you can run the machine.

The federal market is not a lottery. It is a pipeline — and the pipeline is teachable to a person and to a machine. One of them signs.