

Facilitator Guide — Simulation Day: The Portfolio Wargame

A full-day, facilitator-run portfolio simulation: multiple capture organizations compete in a fictional shared market, build a pursuit portfolio, score it, choose what to bid, defend a bid/no-bid gate, and present one submission-readiness review before a judging panel. The single-pursuit simulation of DL 210, raised to portfolio altitude.

What this is. DL 210 runs each student through *one* pursuit for a term. This day runs the *portfolio* a capture organization actually manages: a room of teams, each a fictional company, all competing for the same fictional market. The team that wins is not the one that writes the best single proposal — it is the one whose **portfolio discipline, gate judgment, and readiness presentation** hold up under capacity pressure. Run it as a capstone day for DL 210, a mid-program intensive, a fellowship cohort’s final exercise, or a professional workshop.

The stack rule holds. Every artifact in this pack is produced with paper, a whiteboard, or a spreadsheet. No software required, none tested. The machine side of the doctrine can be a *reference* in the room (see Extensions), but the call is always human.

The hidden teaching device. The facilitator holds a sealed envelope — the fictional market’s *actual outcomes*. The debrief opens it and teams score their calibration. That is the day’s real lesson, and it is built in, not bolted on.

1. The day at a glance

Six blocks. Times are a template for a 9:00–17:15 day; compress per Extensions.

Block	Time	Focus	Deliverable
1 — The market read	9:00–10:30 (90)	Read-in, team formation, market segmentation	One-page segmentation + portfolio thesis
2 — Portfolio construction	10:45–12:00 (75)	Score every opportunity with RICE; first pWin sketch	RICE board (all four opportunities)

Block	Time	Focus	Deliverable
3 — Score, gate, allocate	12:45–14:15 (90)	Full pWin × value scores; bid/no-bid gates; capacity allocation	Portfolio scorecard + gate charters + allocation
4 — The readiness review	14:30–16:00 (90)	One pursuit, presented as ready to submit, defended to the panel	10-minute presentation + defense per team
5 — Panel deliberation	16:00–16:30 (30)	Panel scores against the rubric; awards	Scores + decisions
6 — The debrief	16:30–17:15 (45)	Calibration reveal; portfolio post-mortem; the learning loop	One-page post-mortem per team

Between blocks 2 and 3 sits lunch; the blocks themselves are the spine.

2. Setup and materials

The room. 3–4 tables (one per team), a wall or board per table for the portfolio board, and a presentation space for block 4. A clock that everyone can see — the pipeline has deadlines, and so does the day.

The panel. Two or more judges: faculty, and ideally a guest practitioner (a capture/proposal professional) or a visiting executive. The panel sits through block 3’s gates and block 4’s presentations, and deliberates in block 5. Brief the panel with the panel briefing before the day.

Materials to print, per team:

- The **market pack** (the four fictional opportunities — §4) — one per team, identical.
- The **team kit** for that team’s company (§5) — each team gets only its own.
- A **blank RICE board** — a grid of the four opportunities × the four RICE factors.
- A **gate-charter template** — the four parts (trigger, criteria, decision, owner) per gate.
- A **capacity-budget worksheet** — each opportunity’s pursuit cost vs. the team’s pursuit budget.

- A **submission-readiness checklist** — the pre-submission gate’s criteria (§6, block 4).
- The **scoring rubric** (§7) — so teams know exactly what the panel is judging before the day ends.

The facilitator’s sealed envelope. The fictional actual outcomes (§9) — who wins each opportunity, at what price posture, and the outcome the teams’ scores will be measured against. No one opens it until block 6.

3. The scenario brief — the fictional market

Read to the room (or print). The **Meridian Corridor** is a fictional regional economy anchored by four federal buyers and a growing base of small and mid-size federal contractors. For one day, your team *is* a federal contractor competing in it. Four opportunities are live or imminent across four agencies. Every team in this room is a competitor for the same market. Your job is to run your company the way a real capture organization runs itself: segment the market, score the opportunities, decide what to bid, gate your choices honestly, and stand behind the one pursuit you present as ready to submit. You have finite people and a finite pursuit budget. Chasing everything is the one strategy guaranteed to fail.

The scenario’s four opportunities are in §4, and the four companies are in §5. The market is deliberately shaped so that **no single team can win everything** — each opportunity favors a different posture, and the capacity math forces each team to give something up.

4. The market data — four fictional opportunities

All four opportunities are shared; every team sees the same pack. The “signal” column is what a sharp capture team reads between the lines — it is not printed on the opportunity fact sheet, it is what the team is expected to notice.

Opportunity A — NGRA: GridSense Analytics Support Services

Fact	Value
Agency	National Grid Resilience Agency (NGRA) — a fictional federal agency responsible for grid-modernization analytics
Vehicle	Single-award services IDIQ; 1 base year + 3 option years

Fact	Value
Value	~\$18M all-options
Set-aside	HUBZone set-aside (NAICS 541512)
Incumbent	Apex Utility Systems (fictional large prime) — 9 years in place; strong customer relationship; visible pricing creep and a stale toolchain
Evaluation	Technical 45 · Management 25 · Past Performance 20 · Price 10
Key requirements	Named program manager + analytics lead; 30% small-business subcontracting plan; a transition plan; a published grid-resilience metric as a quantified outcome target
The signal	A recompete of an incumbent whose pricing crept. The agency published a “lessons learned” memo valuing innovation over continuity — a displacement window for a HUBZone firm with a better analytic story.

Opportunity B — DSCI: Supply Chain Intelligence Data Platform

Fact	Value
Agency	Department of Homeland Security — Directorate for Supply Chain Intelligence (DSCI) — a fictional directorate
Vehicle	Single-award FFP platform build + 2-year sustain; FAR Part 15
Value	~\$25M all-options
Set-aside	Small-business set-aside (NAICS 541519)
Incumbent	None — white space ; the agency has been prototyping internally with a small prime contractor (the shadow incumbent)
Evaluation	Technical 50 · Management 20 · Past Performance 15 · Price 15

Fact	Value
Key requirements	A modern integration architecture; a security-clearance posture for a small team; an AI-assisted data-integration capability; a day-one operating capability — the agency is under a congressional deadline
The signal	White space, but a hard deadline. Speed and a concrete security posture are the discriminators. The internal prototype firm is the shadow incumbent and must be displaced or beat on technical depth.

Opportunity C — NPHDI: EpiShare Interoperability Program

Fact	Value
Agency	National Public Health Data Institute (NPHDI) — a fictional institute within Health and Human Services
Vehicle	Cost-plus services; 1 base year + 2 option years
Value	~\$12M all-options
Set-aside	WOSB set-aside (NAICS 541690)
Incumbent	DataBridge Partners (fictional small incumbent) — well-liked by program staff, but the agency is publicly unhappy with cost overruns and thin capacity
Evaluation	Technical 40 · Management 25 · Past Performance 20 · Price 15 (cost realism applies)
Key requirements	Data-interoperability standards work (the standards named generically, not as a product); a program manager with public-health domain credibility; a realistic, auditable cost buildup

Fact	Value
The signal	A loved-but-strained incumbent. The agency's own inspector-general flagged cost controls — disciplined cost posture and program-management credibility are what displaces DataBridge.

Opportunity D — RRA: Autonomous Resilience Modeling, Phase I

Fact	Value
Agency	Resilience Research Agency (RRA) — a fictional federal advanced-research agency
Vehicle	STTR Phase I research award (~6 months) with a Phase II path
Value	~\$250K Phase I; Phase II ~\$1.5M
Set-aside	STTR — requires a small business teaming with a nonprofit research institution
Incumbent	None (open topic)
Evaluation	Technical merit and innovation · the research team and institution · commercial potential · the Phase II path (judged by topic reviewers)
Key requirements	A named research-institution partner; a compelling technical abstract; a small-business principal investigator; eligibility certifications
The signal	The low-cost option-value bet — a cheap entry that seeds a research lane and a Phase II follow-on. The discriminator is the strength of the research partnership, not price.

The portfolio shape. A is a HUBZone recompetete play; B is a white-space platform bet; C is a WOSB displacement play; D is a cheap research option. The teams that read the set correctly notice: the set-asides split the market (A and C each favor a specific certification), B rewards speed and security depth, and D rewards a research partnership only two of the four companies can honestly claim.

5. The team kits — four fictional capture organizations

Each team gets only its own kit. The kits are balanced: every company can win *something* and must give up *something*. The liabilities are the point — they are what make the bid/no-bid gate real.

Kit 1 — Meridian Analytics Group

Profile	A ~\$42M revenue small business, ~130 staff, headquartered in a mid-Meridian-corridor city; data analytics, applied modeling, cloud engineering
Certifications	HUBZone ; SAM-registered; NAICS 541512, 541330, 518210
Partner	The Alleghany Institute for Applied Research (fictional regional research university) — makes Meridian STTR-eligible
Past performance	Strong state and federal analytics work in energy and transportation; thin on federal platform builds
Capacity	Can staff ~2 concurrent pursuits; pursuit budget ~\$60K
Posture	The regional HUBZone player — wins where the set-aside, the local story, and applied analytics align
Liabilities	Not a proven platform builder; civilian-agency relationships outside the region are thin; a shallow program-management bench

Kit 2 — Ferro Engineering Group

Profile	A ~\$36M revenue woman-owned small business, ~95 staff; systems engineering, program management, civil/infrastructure
Certifications	WOSB + EDWOSB ; NAICS 541330, 541690
Partner	None permanent; buys niche subcontractors as needed

Past performance	Disciplined program/project management across state and federal infrastructure programs; strong performance ratings
Capacity	Can staff ~2 concurrent pursuits; pursuit budget ~\$45K
Posture	The disciplined program manager — wins on management quality, cost credibility, and set-aside access
Liabilities	A thin data-science/ML bench; has never run a platform integration at DSCI's scale; no research lane

Kit 3 — Lumen Peak Solutions

Profile	A ~\$20M revenue small business, ~60 staff, spun out of a national-laboratory network; physics, machine learning, resilience modeling
Certifications	SAM-registered; NAICS 541715, 541330; no set-aside certifications
Partner	An exclusive partnership with the fictional Cascade National Laboratory — deep research credibility, STTR-active
Past performance	SBIR/STTR wins and lab-origin prototypes; almost no civilian-agency services experience
Capacity	Can staff ~1–2 concurrent pursuits; pursuit budget ~\$35K
Posture	The technologist — wins on research depth, white-space integration, and the lab partnership
Liabilities	Weak customer relationships in civilian agencies; a small bench for a services engagement; no past performance as a prime on services contracts

Kit 4 — Northline Federal Services

Profile	A ~\$250M revenue mid-size prime, ~700 staff; broad federal services across defense and civilian agencies
Certifications	Full-and-open (no small-business set-aside eligibility); NAICS 541330, 541512, 541519, 518210
Partner	A mature subcontractor network; can staff anything
Past performance	Deep prime record; strong but expensive (high overhead)
Capacity	Can staff ~3–4 concurrent pursuits; pursuit budget ~\$150K
Posture	The prime — competes everywhere, wins on scale, breadth, and delivery depth
Liabilities	High overhead makes the small set-asides uneconomic (a winner’s-curse hazard on price); slow; hard to tell a small-business story where a set-aside matters

6. The six blocks — facilitator scripts

The scripts are the facilitator’s **speaking notes**: what to open with, what to plant, what to watch for. Keep the times; they are the discipline.

Block 1 — The market read (90 min)

Open (5 min). “For one day, this room is a federal market. You are not students; you are the companies in your kits, competing for the opportunities in the pack. The day is not about the best single proposal — it is about the best portfolio under pressure. Read your kit, read the market, and decide where your company plays.”

Do now (60 min). Each team reads its kit and the market pack, then produces a **one-page market segmentation + portfolio thesis**: which lanes the company will play (by set-aside, by agency, by capability), which it will not play, and why. The thesis is the answer to the doctrine’s first gate question — *where do we play?* — written before any opportunity is chased. The facilitator’s planted question at 30 minutes: **“Which lane is yours and not just available?”** A team that names every lane has not segmented; it has listed.

Report out (20 min). Two minutes per team: the thesis, one sentence. The panel hears the portfolio posture of every team before any score exists — this is

the baseline the rest of the day is judged against.

Close (5 min). “Segmentation is a decision, not a description. If your thesis survives into the afternoon unchanged, you are either brilliant or you did not read the market. Keep it honest.”

Watch for: teams that refuse to give up any lane (they will choke on capacity later); teams that copy the kit’s posture without reasoning about the market’s signals.

Block 2 — Portfolio construction with RICE (75 min)

Open (10 min). “You have a capacity constraint and four opportunities. Before you score them in depth, you need an order. The portfolio tool is RICE — **Reach** $\times$ **Impact** $\times$ **Confidence** $\div$ **Effort**. Reach is how much the opportunity matters; Impact is how much a win changes your company; Confidence is how sure you are you can win it; Effort is what pursuing it costs. Score all four opportunities on each factor, compute the RICE number, and rank them. RICE tells you the order to *look at*; it does not make the bid/no-bid call.”

Do now (50 min). Teams build the RICE board. The facilitator’s planted questions: “**What is ‘reach’ for a set-aside you are not certified for?**” (answer: zero — that is a hard knockout showing up in the confidence factor). “**Is RRA’s low dollar value low reach, or low effort with option value?**” (the sharp teams score effort and Phase II option value, not just the Phase I dollar).

Board review (15 min). Each team posts its ranked board. The facilitator asks one question per team: “**Which ranking would your company’s competitor hope you got wrong?**” — the day’s first test of whether the ranking is an argument or a spreadsheet.

Watch for: teams that rank purely by dollar value (they miss D’s option value and B’s strategic value); teams whose confidence factor never falls below 80% (overconfidence is the disease the whole day exists to cure).

Block 3 — Score, gate, allocate (90 min)

Open (10 min). “Now you do the real work: full scores, hard gates, and an allocation that fits your capacity. For each of the four opportunities, produce a **pWin**, argued factor-by-factor with evidence from the market pack, a **value** grounded in the solicitation’s numbers, and a composite **pWin** $\times$ **value** against a **threshold** your company sets in advance. Then run the **bid/no-bid gate** — trigger, criteria written before the work, a decision (BID / BID WITH CONDITIONS / NO BID), and a named owner. Finally, fit your allocation inside your pursuit budget. The pursuit cost per opportunity is in your pack: A ~\$50K, B ~\$75K, C ~\$40K, D ~\$15K. You cannot bid everything.”

Do now (60 min). Teams produce the portfolio scorecard, the gate charters,

and the capacity-budget worksheet. The facilitator’s planted questions: “**Which opportunity is your disciplined no — and who owns it?**” “**If your pWin says NO BID but your strategy says BID, where is that override recorded, and who is accountable?**” (the doctrine’s threshold-override-as-governance-decision, not a silent exception).

Gate report (20 min). Each team states its four dispositions in one minute each. The panel listens for the shape: a team that bids everything has not gated; a team with at least one hard NO BID is showing judgment. The panel may ask one question per team: “**What would have to be true for that NO BID to flip?**”

Watch for: the winner’s-curse bid (Northline chasing a set-aside it cannot price); the not-certified bid (Meridian bidding B despite no platform proof; Ferro bidding B despite no integration bench); the fear-of-missing-out bid (Lumen bidding C despite no services past performance). Each is a gate that should have said no.

Block 4 — The readiness review (90 min)

Open (5 min). “The portfolio is chosen. Now you stand behind one pursuit as if it must ship: the **submission-readiness review** — the last gate that can stop a bad submission. Present the pursuit your team is readiest to submit as if you are asking for executive sign-off. The panel will judge the presentation; the rubric is in your pack.”

Present (60 min). Each team presents **10 minutes**, then defends **5 minutes** against the panel. The readiness review must show: the win strategy (who the customer is, what they are measured on, the discriminator), the teaming direction (who is in and why), the key personnel (named, credible), and the **pWin defended** — the same estimate from block 3, now argued as a submission-ready claim. The panel’s stance is the capstone’s: find the gate that would fail first and ask why it is not already written.

Score (25 min). The panel scores each team against the rubric (§7) while the room watches the clock. The scores feed block 5.

Watch for: teams that present the pursuit they are proudest of instead of the one their portfolio logic says is readiest; teams that cannot name the residual risks (a readiness review that is all confidence is not a review, it is a pitch).

Block 5 — Panel deliberation (30 min)

What happens. The panel reconciles scores against the rubric, names the winners, and the facilitator announces: **Best Portfolio Discipline, Best Gate Judgment, Best Readiness Presentation**, and the **Overall Winner**. The deliberation is short on purpose — the day’s real ending is the debrief, not the trophy.

Block 6 — The debrief protocol (45 min)

The debrief is the day’s engine of learning, and it runs on the sealed envelope. Protocol:

1. The calibration reveal (10 min). The facilitator opens the envelope and reads the fictional actual outcomes (§9): which company won each opportunity, at what price posture, and what the agency said about why. Teams then compute their calibration score: of their BID dispositions, how many landed on an actual winner? The point is not shame — it is the doctrine’s own discipline, that a pWin estimate must be *true* (an agent that says 0.7 should win about seven times in ten).

2. The portfolio post-mortem (15 min). Each team works its one-page post-mortem: for each opportunity, what did we decide, what happened, and — with the outcome known — would we decide the same? The questions that matter: “*Where did we over-score ourselves?*” “*Which NO BID turned out to have cost us a real win — and was that the right call anyway?*” “*Which BID were we lucky on?*” The post-mortem is graded on honest comparison, not on being right — exactly the standard Doctrine 10 sets.

3. The learning loop (10 min). Teams report one line each: the single factor their calibration most over-weighted, and the one change they would make to next cycle’s pWin factors. The facilitator’s close: “**You just ran the learning loop. Every loss here is data, and data compounds. The organization that gates honestly and debriefs without flinching is the one that wins the next cycle — and that is the whole point of this day.**”

4. The individual reflection (10 min). Silent, written: “*What did I learn about my own over-optimism?*” — the doctrine/04 lesson that humans systematically overestimate the value of the thing in front of them. The day ends with that paper in hand.

7. The scoring rubric

The panel scores every team on three dimensions — **portfolio discipline, gate judgment, and presentation**. Four-point scale, matching the course’s convention (Not Ready · Developing · Proficient · Distinguished).

Dimension	What the panel is judging	Not Ready	Proficient	Distinguished
Portfolio discipline (40%)	Did the team segment the market, score every opportunity (RICE + pWin), respect its capacity, and build a portfolio that is a strategy and not a wish list?	Chased every opportunity; no segmentation; capacity ignored; the portfolio is a folder	Segmented; scored; the allocation fits the pursuit budget	The portfolio has a thesis — a lane chosen and defended; the allocation is argued; the team names what it gave up and why
Gate judgment (30%)	Were the bid/no-bid dispositions defensible — criteria written in advance, hard knockouts honored, at least one disciplined “no”?	Calls asserted; criteria invented at the meeting; bid on everything	Criteria written before the work; dispositions defensible; a “no” appears somewhere	The hard “no” is named and owned; BID WITH CONDITIONS is governed, not hand-waved; any threshold override is a recorded decision with a named owner
Presentation & readiness (30%)	Was the chosen pursuit presented as ready to submit — win strategy clear, customer understood, pWin defended, survives the panel?	Vague; pWin asserted; evasive under questioning	Clear; customer map shown; pWin argued factor-by-factor; defends the call	The readiness review names residual risks before the panel finds them; the pitch survives a hostile panel and the team knows its own limits

Composite: Portfolio discipline 40 · Gate judgment 30 · Presentation & readiness 30. The portfolio gets the largest weight because that is the altitude this day adds over DL 210's single pursuit.

The trap being tested. The single-pursuit simulation rewards the best proposal. This day rewards the **best portfolio** — and the team that treats the day as four separate proposal competitions loses, because capacity is the scarce resource and the gate is the discipline that protects it. The panel's sharpest questions are the ones that probe *why this pursuit, in this portfolio, over the alternatives*.

8. The panel briefing

Give the panel this page before the day. The stance is the capstone's: **skeptical, and looking for the gate that would fail first.**

- **Portfolio discipline is judged early.** Listen to block 1's thesis and block 3's dispositions as one continuous argument. A team whose afternoon portfolio contradicts its morning thesis is showing the failure — note it.
 - **Gate judgment is judged on the “no.”** The panel should positively reward the team that walks away from an opportunity the market made attractive but their kit made uneconomic. A portfolio with no NO BID is a portfolio with no gate.
 - **Presentation is judged on the defense.** The 10 minutes are the argument; the 5 minutes of defense are the evidence. Ask the question that forces the team to name a limit: *“What in this package would a rival's black-hat review cut first?” “What is the one factor your pWin is softest on, and what is your plan if it is worse than you scored?”*
 - **Score against the rubric, not against each other.** Each team is judged on the dimensions, not on whether it beat the other teams. The winners emerge from the rubric, not from the panel's memory of the loudest table.
 - **The day's real winner is not the trophy.** The panel's written comments should feed each team's post-mortem. The sealed envelope is the teacher; the panel is the referee.
-

9. The sealed envelope — fictional actual outcomes

The facilitator opens this in block 6. The outcomes are designed so that a well-run team scores well and an overconfident team scores poorly — the calibration lesson lands either way. Print this page separately and seal it.

The fictional market's actual outcomes:

- **Opportunity A (NGRA — HUBZone recompetete):** won by **Meridian Analytics** at a price modestly below Apex’s incumbent price, on the strength of the quantified grid-resilience outcome target and the transition plan. Apex’s pricing creep and the agency’s “innovation over continuity” memo were the deciding factors. *A sharp Meridian plays this and wins; a Meridian that spread itself across A, B, and C loses capacity and still wins A, but weakly — its bid/no-bid judgment is what the panel scores.*
- **Opportunity B (DSCI — white-space platform):** won by **Lumen Peak Solutions** on the technical factor — the Cascade lab partnership and the day-one capability beat the shadow incumbent on integration depth, and the agency accepted a realistic (not lowest) price for the security posture. *This is the white-space bet that rewards the technologist. A Lumen that instead chased A for its own certification advantage misread its own kit.*
- **Opportunity C (NPHDI — WOSB displacement):** won by **Ferro Engineering Group** on management and cost realism — the auditable cost buildup and the program-management credibility displaced DataBridge, whose cost overruns the IG had flagged. *This is the disciplined program manager’s lane. A Ferro that bid it as a price shootout instead of a cost-credibility play lost it.*
- **Opportunity D (RRA — STTR Phase I):** won by the team with the strongest research-institution partnership and the most credible Phase II path — for a Meridian with the Alleghany partnership or a Lumen with Cascade, it is a cheap win; for a team without a research partner, **it was never a real opportunity**, and bidding it was the tell of a team that had not read its own kit.

The calibration scoring: award each team a point for every BID that landed on a winner, and — just as important — note whether each team’s NO BID or BID WITH CONDITIONS on the opportunities it declined *correctly predicted the winner’s identity*. A team that correctly says “we should not bid D because we have no research partner” shows better judgment than a team that bids D and loses, even though both lose D.

10. Extensions and variants

- **The compressed half-day (blocks 1–3 + a short report-out).** Run the market read, the RICE board, and the gate dispositions, and skip the full readiness presentation. The day becomes a portfolio-strategy workshop: the deliverable is the scorecard and the gate charters, and the debrief is

shortened to the calibration reveal.

- **The two-day version.** Add a full color-team block between blocks 3 and 4: each team runs pink (compliance) and red (evaluator-stance) on its chosen pursuit, with a second room's team as the reviewers. The readiness presentation then defends against the color-team findings. This is the DL 302 + portfolio day combined.
- **The machine-assist variant.** Permit teams to consult an AI layer for a second opinion on a RICE score or a pWin factor — but the call is always human, and the panel may ask a team why it overrode or accepted the machine's read. This is the doctrine/08 posture made live: the tool is an accent, the concept is the language.
- **The surprise-drop variant.** At the start of block 3, the facilitator drops a fifth opportunity — a small, set-aside-matched, short-fuse award — and teams must reallocate within 15 minutes. It tests whether the portfolio thesis survives a live surprise, which is the entire point of Doctrine 04's "the further a pursuit advances, the more expensive the decision to stop."

11. Where this sits in the curriculum

- **It extends DL 210.** The single-pursuit simulation (`./course/undergraduate-210-capture-and-pursuit`) teaches one pursuit end to end. This day teaches the *portfolio* the doctrine's gates protect — the same tools (RICE + pWin), the same gates, raised from one pursuit to many, under a shared-market competitive frame.
- **It exercises the doctrine's core.** Segmentation and gates (`./doctrine/04-gates-and-governance.md`); scoring, thresholds, and price-to-win (`./doctrine/05-scoring-and-price-to-win.md`); capture strategy and discriminators (`./doctrine/09-capture-and-competitive-strategy.md`); the debrief and the learning loop (`./doctrine/10-post-submission-and-win.md`).
- **It uses the portfolio vocabulary.** RICE as the prioritization scorecard, and calibration as the discipline that a pWin must be true — both taught in `./literacy/how-the-machine-learns.md`. The shared-core scoring exercise (`./modules/shared-core.md`, S5) is this day's engine at single-opportunity scale.
- **It bridges to the graduate tier.** The portfolio frame is exactly what the graduate capture capstone runs at firm scale (`./practice/capstones/graduate-capture-capstone.md`) — a mid-size firm's capture manager deciding what to bid as a business. This day is the undergraduate and fellowship door into that altitude.
- **The rubric echoes the course's contract.** The four-point scale, the public rubric, and the "criteria written before the work" discipline come from `./course/assessments-and-rubric.md`; the panel's stance comes from the capstone defense.

The day in one sentence

A room of competing capture organizations, a shared fictional market, and one day to prove the portfolio discipline the gates protect: segment, score, gate, allocate, present — and then open the envelope and let the real outcomes teach you what your pWin was actually worth.