

Doctoral Advanced Module — Advanced Capture Management as a Research Domain

The master's tier teaches you to run a capture. This seminar treats capture itself as the research object — the bid/no-bid decision as a decision problem, the capture organization as a unit of analysis, the pWin estimate as an instrument to calibrate, and price-to-win as a construct to measure.

Level: doctoral (L3 Architect / L4 Networker) · **Format:** one 3-hour seminar per week, five sessions · **Prerequisites:** DL 901 (design and measurement); DL 902 (statistics and causal inference) recommended · **Companions:** doctrine/04, doctrine/05, doctrine/09, modules/doctoral/research-agenda.md · **Stack-agnostic promise:** capture management is studied as a set of durable decisions and organizations — never as a product's workflow (doctrine/08).

This module is the second of the three C.A.S.E.-adjacent advanced modules the degree crosswalk promises at the doctoral tier (alongside AI systems and partnerships-and-ecosystems). Where DL 903 studies the discipline's *empirical machinery*, this seminar studies the **craft itself as a research domain**: the decisions capture managers make, the organizations that make them, and the constructs — pWin, price-to-win, gate discipline — that the field argues about without evidence. Almost every one of these arguments is researchable.

Course description

Capture management is practiced at scale — thousands of firms, billions of dollars — and argued about constantly: *should we have bid? which factor mattered most? is our pWin honest? did we price to win or price to lose? does our gate discipline pay?* The field's tragedy is that nearly all of these arguments are settled by seniority and anecdote rather than by evidence. This seminar exists to convert them into research questions, study designs, and defensible contributions.

The seminar moves up four levels of analysis, one per block:

1. **The decision.** The bid/no-bid call is the discipline's most consequential decision — a decision under uncertainty that decision theory can formalize and the survivorship hazard constrains.
2. **The organization.** The capture organization — its gates, charters, decision rights, and pipeline — is a coherent unit of analysis, comparable across firms.
3. **The estimate.** pWin is a measurable estimate, and its honesty — calibration, reliability, factor validity — is testable. This is the calibration research program.
4. **The price.** Price-to-win is a position chosen under competition, and it can be studied with the tools of cost research and auction theory.

The deliverable is a **seminar research design** — one study, fully designed, defensible to a skeptical panel — that sits on the research agenda and is designed to grow into the dissertation’s empirical core.

Learning outcomes

By the end of this module, a student can:

1. **Formalize** the bid/no-bid decision as a decision under uncertainty, and name the sampling hazard that constrains every study of it.
 2. **Treat** the capture organization as a unit of analysis — define its gate structure, decision rights, and pipeline metrics as measurable variables.
 3. **Design** a calibration and factor-validation study of the pWin estimate, using DL 902’s statistical toolkit.
 4. **Operationalize** price-to-win as a research construct — cost floors, competitive ranges, margin outcomes — and connect it to auction theory.
 5. **Defend** a seminar research design that a dissertation committee would accept.
-

Session 1 — Bid/no-bid as a decision problem

The claim. The bid/no-bid call is the discipline’s most consequential decision, and it is a decision under uncertainty: a bet placed on an estimate of winning, with an expected value the field could compute — but usually does not.

The ideas. The decision can be formalized: an organization pursues when the expected value (probability of win $\times$ value, less the cost of pursuit) clears a threshold (doctrine/05). Decision theory supplies the vocabulary — subjective probability, expected value, loss functions, the utility of a pursuit beyond its dollars (past-performance compounding, market entry). But the decision is made by an organization under a published rule structure, which is what makes it *researchable* rather than merely philosophical.

The central hazard, named early. The research constraint that shapes every study in this module: **only pursued bids resolve**. You can observe what happened to the pursuits an organization chose to pursue; you cannot observe what would have happened to the opportunities it declined. The sample of resolved outcomes is chosen by the very decision under study. The seminar treats this not as a footnote but as the design constraint that separates serious capture research from naive capture research.

The hidden assumptions. That the pWin estimate and the value estimate are honest; that the threshold is set correctly; that the survivorship hazard can be mitigated (observation floors, stratification, a stated withdrawn-as-non-win policy); that “should we have bid?” is even answerable from the record.

The research questions. This is the decision-and-governance thread (Thread D):

- Does **threshold discipline** improve portfolio outcomes — do firms with honest bid/no-bid gates convert and win at better rates than firms that chase everything?
- Does the **bid/no-bid call itself** discriminate — do declined opportunities really resolve worse than pursued ones? (The survivorship hazard is the whole ballgame.)
- Is the decision **well calibrated at the portfolio level** — does the firm’s aggregate expected value match its aggregate realized value?

Design exercise. Teams design the **threshold-discipline study**: the data (firms’ scored pursuit sheets, bid/no-bid decisions, and terminal outcomes over a defined window — organizational data, ethically gathered), the design (comparative case or difference-in-differences around a governance change), and the mitigation of the survivorship hazard.

Reading. doctrine/03 (the pipeline — where the decision sits), doctrine/04 (gates), doctrine/05 (the score and threshold); research-agenda Thread D.

Session 2 — The capture organization as a unit of analysis

The claim. The capture organization is not just a team; it is a decision structure — gates with written criteria, named owners, decision rights, and pipeline metrics — and it is a coherent unit of analysis that can be described, compared, and studied.

The ideas. An organization’s capture system is its *operating system for pursuing work*: the gate ladder (how many gates, what criteria, who owns each), the engagement cadence (how the team interacts with the customer and the market), the bid/no-bid authority (who can say no), the pipeline metrics (what the organization measures about its own pursuit), and the review discipline (how color teams are run). The discipline’s doctrine encodes the design ideal (doctrine/04, doctrine/09); the empirical question is what actually predicts outcomes.

Treating the capture organization as a unit of analysis opens the comparative case design: two or three firms (or two eras of the same firm) compared on gate structure, decision rights, and outcomes. The construct of “gate discipline” can be operationalized many ways — share of scored pursuits declined below threshold, existence of written gate criteria before pursuit start, the escalation path when a gate is overridden — and the measurement choice is itself substantive (DL 901 Unit 4).

The hidden assumptions. That the organization’s written structure matches its actual decision behavior (the gate charter that is never enforced); that gate discipline is a stable trait rather than a response to circumstance; that comparable firms are comparable.

The research questions. Thread D again, at organizational altitude:

- Does **gate governance** — criteria written before the work, named decision owners — improve the quality of pursuit decisions and the honesty of the pipeline?
- What distinguishes firms that **say no** from firms that chase everything — and does the discipline pay?
- How do **decision rights** (who owns the bid/no-bid call, the gate, the submission decision) shape a portfolio?

Design exercise. Teams design the **comparative-case study of capture governance**: the firms, the measurement protocol (a gate-governance codebook scored by independent coders with inter-rater reliability), the outcome measures (conversion, win rate, portfolio value), and the rival explanations the design must rule out.

Reading. doctrine/04 (gates and governance) in full; doctrine/09 (capture as competitive strategy); DL 901 Unit 4 (measurement) and Unit 3 (comparative case designs); the graduate organizational-theory course (course/graduate-organizational-theory-governance.md) at altitude.

Session 3 — The pWin calibration research program

The claim. pWin is the field’s central estimate — the probability of win that drives the score, the threshold, and the bid/no-bid call — and its honesty is testable. A 0.7 estimate should win roughly 70% of the time.

The ideas. The discipline models pWin as a seven-factor composite — past performance, technical capability, price competitiveness, incumbent/relationship advantage, team strength, customer knowledge, strategic alignment — each scored on a scale and weighted into a 0-to-1 estimate (doctrine/05). For a researcher, that composite is a **measurement instrument** with three properties worth studying:

- **Validity** — does the composite predict awards, and do its factors and weights match what the data show?
- **Reliability** — do two experienced scorers score the same opportunity similarly? (If not, the composite is measuring the scorer, not the opportunity.)
- **Calibration** — does a 0.7 mean roughly 70%? (Reliability diagrams, Brier score, and expected calibration error localize the answer.)

The calibration research program is the field’s most fundable open question because the *data exist*: the prediction-and-outcome record stores, for each pursued opportunity, the predicted pWin at decision time and the terminal outcome. The pairing is the discipline’s empirical goldmine.

The hidden assumptions. That factor scores are assigned reliably; that a single weight set serves all agencies and funder families; that the observation floor is adequate; that “withdrawn” is a non-win.

The research questions. Threads A and B:

- Is the composite **calibrated** — and in which bands is it over- or under-confident?
- Which **factor weights** actually move the probability of winning — and do the discipline’s hand-set weights match the data?
- Do the factors that predict wins at one agency hold at **another** — or is factor validity agency-specific?
- Do **structured composites** calibrate better than **expert judgment** — the field-experiment question?

Design exercise. Teams design the **calibration and factor-validation study**: the data (predicted factor scores + pWin + outcomes for a sample of pursued bids), the model (logistic regression of won/lost on the seven factors), the calibration statistics (reliability bins, Brier, ECE), the reliability protocol (two coders scoring a shared set of opportunities), and the preregistered decision rule for “this composite is miscalibrated.”

Reading. doctrine/05 in full; DL 901 Unit 4 (measurement) and Unit 5 (sampling and the survivorship hazard); DL 902 Movement 2 (logistic regression) and Movement 3 (calibration); research-agenda Threads A and B.

Session 4 — Price-to-win as a research construct

The claim. Price-to-win is not a guess; it is a position chosen under competition — and as a research construct it connects the discipline to cost research and auction theory in ways the field has barely exploited.

The ideas. The discipline teaches price as a position: above the cost floor (what it costs us to do the work), inside the competitive range (what the field will land in), and set to the evaluation (what the agency rewards — price or best value) (doctrine/05). As a research construct, that position has measurable parts:

- **The cost floor** — the honest cost of performance, the margin floor per contract type. Researchable as an accounting-and-pricing question: what do firms actually bid relative to their floors?
- **The competitive range** — the estimated band of competitor prices. Researchable as a distributional question: how wide is the range, and how accurately do firms estimate it?
- **The margin outcome** — what the winning bid actually yields. Researchable as an auction-theory question.

Auction theory supplies the discipline’s sharpest hypotheses: federal competitions are, at heart, first-price auctions with a published rulebook and an observant rival. The theory predicts **bid shading** (bidding below true value to win), the **winner’s curse** (the winner is the bidder who most overestimated value, so winning itself is bad news), and the strategic logic of when an aggressive posture is rational. The research question is whether the theory’s predictions show up in realized margins and win/loss records.

The hidden assumptions. That firms actually compute floors and ranges rather than anchoring on competitors; that the competitive range is knowable; that the winner's curse operates in a market with published prices; that an aggressive posture that wins is a win rather than the start of a loss.

The research questions.

- Does **bid shading** show up in realized margins — do winning bids sit below the value of the work, and does that predict post-award distress?
- Does the **winner's curse** operate in federal competitions — do the most aggressive winning bids predict the worst delivery outcomes?
- Does an **aggressive pricing posture** (price-to-win at the floor) win more but lose more on delivery — or is the cost-technical trade real?
- How well do firms **estimate the competitive range**, and does estimation accuracy predict outcomes?

Design exercise. Teams design the **price-to-win study**: the data (winning and losing bid prices, the public award values, and — where obtainable — internal floors and range estimates), the design (an observational study of bid behavior and margins, or a field experiment on pricing-posture training), the analysis (margin outcomes by posture, winner's-curse tests), and the honesty hazards (selection into aggressive posture).

Reading. doctrine/05 (price-to-win as a position); course/graduate-microeconomics-of-procurement.md and course/graduate-elective-negotiations-pricing-posture.md at altitude; the auction-theory canon — common-value auctions, bid shading, the winner's curse — as named concepts; the debrief and protest record as the post-award evidence base.

Session 5 — Synthesizing the research program: design workshop

The final session is a working session and a defense. The four levels of analysis — the decision, the organization, the estimate, the price — are one story: the discipline *decides* (Session 1) inside an *organization* (Session 2) using an *estimate* (Session 3) and a *price* (Session 4), and it claims to *learn* from the outcomes. Each student presents their seminar research design — one study that tests one claim at one level of analysis — and the room attacks it the way a dissertation committee would: *What is the identifying assumption? Where is the survivorship hazard? Is the sample big enough? Who scores the factors, and how do you know they agree? What exactly will you claim if the data go the other way?*

The panel's job is the committee's job: push until the design's limits are named, then judge whether the design is *defensible*, not perfect.

Research-question bank

Each question is dissertation-sized: researchable, feasible, and a contribution the field needs.

1. **Threshold discipline.** Does honest bid/no-bid gating — declining below-threshold pursuits, writing criteria in advance — improve conversion, win rate, and portfolio value?
2. **The bid/no-bid call as a discriminator.** Do declined opportunities really resolve worse than pursued ones — and how would you know, given that only pursued bids resolve?
3. **Gate governance.** Does an organization's gate structure — count, criteria, ownership, escalation path — predict the quality of its pursuit decisions?
4. **Decision rights.** Who owns the bid/no-bid call and the submission decision, and does that ownership structure predict portfolio outcomes?
5. **Factor validity.** Which of the seven pWin factors actually move the probability of winning, and do the discipline's hand-set weights match the data?
6. **Cross-agency stability.** Do the factors that predict wins at one agency or funder family hold at another — or is a single weight set a fiction?
7. **Reliability of pWin scoring.** Do two experienced capture managers score the same opportunity similarly — or is the estimate measuring the scorer?
8. **Composite vs. expert.** Do structured composites calibrate better than expert judgment, and under what conditions?
9. **Bid shading.** Do winning bids sit below the value of the work, and does aggressive bid shading predict post-award distress?
10. **The winner's curse.** Do the most aggressive winning bids predict the worst delivery outcomes — and does an aggressive posture that wins become a loss?
11. **Portfolio calibration.** Is a firm's aggregate expected value honest — does the portfolio's predicted win rate match its realized win rate?
12. **Learning from losses.** Do organizations that measure outcomes and mine debriefs actually improve their estimates, their gates, and their prices over time?

Reading list

The module is concept-first; the reading is named concepts and schools, not pirated material:

- **The doctrine as the discipline's claims:** doctrine/03 (the pipeline), doctrine/04 (gates and governance), doctrine/05 (scoring and price-to-win), doctrine/09 (capture and competitive strategy), doctrine/10 (post-submission and win — the debrief record).
- **The methods core:** DL 901 (design and measurement), DL 902 (statistics and causal inference), DL 903 (the machinery as a research object) — the

toolkit this module assumes and points at the craft.

- **Decision theory:** expected value and subjective probability; the judgment-under-uncertainty research program (Kahneman and Tversky); the reference-class forecasting tradition. The bid/no-bid call as a decision under uncertainty.
- **Auction theory:** first-price and common-value auctions; bid shading; the winner's curse; the strategic logic of competition with a published rulebook (the auction-theory canon as named concepts).
- **Organizational research:** the organizational-learning literature (the learning-loop concepts of Argyris and the double-loop frame); governance and decision-rights concepts; comparative case method.
- **The graduate canon at altitude:** the microeconomics of procurement, competitive strategy and game theory, and negotiations-and-pricing-posture courses (course/graduate-*.md) — the master's concepts this seminar turns into research constructs.
- **The public record as evidence:** the award record, the solicitation archive, and the debrief and protest record — the ground truth the field argues about.

Dissertation tie-in

This module is the natural home of dissertations on Thread D (decision and governance) and Thread E (learning and evolution), with direct bridges into Thread A (calibration) and Thread B (factor validation). A dissertation from this module typically lands on one: a threshold-discipline study (does gating pay?), a governance comparative case (do decision rights matter?), a calibration study (is pWin honest?), or a price-to-win study (does the winner's curse operate here?). The session designs assemble directly into a dissertation proposal — the four levels of analysis map onto four candidate chapters.

The seminar in two sentences

Capture management is practiced at scale and argued about without evidence — every claim in the craft is a research question waiting for a design. This seminar trains the researchers who will answer them: the decision formalized, the organization measured, the estimate calibrated, and the price tested against the market. **Test the craft so the craft can be trusted.**

Reflection questions

1. Which claim of the craft do you most suspect is *not* true — threshold discipline, factor weights, the winner's curse — and what would your study show if your suspicion were right?
2. For the design you are building, what is the single most likely way a skeptic could dismiss the result, and what have you done to pre-empt it?

3. If a capture executive reads your published study, what is the one sentence of guidance they should be able to take away — and is your design strong enough to support it?