
The Ontology of Operations Management

with 1,000 Applied Diagnostics
across 100 U.S. SMB Industries

Kamyar Shah

2026

Contents

Part I — The Ontology of Operations Management 6

The Ontology of Operations Management 6

PART 0 — FORMAL FOUNDATIONS	6
PART 1 — STRATEGY LAYER	7
PART 2 — DEMAND LAYER	8
PART 3 — PROCESS CORE	9
PART 4 — RESOURCES	10
PART 5 — INVENTORY	11
PART 6 — QUALITY MANAGEMENT	11
PART 7 — PLANNING & CONTROL HIERARCHY	12
PART 8 — SERVICE OPERATIONS	12
PART 9 — PROJECT MANAGEMENT	13
PART 10 — FACILITY & LAYOUT DESIGN	13
PART 11 — SUPPLY CHAIN & SOURCING	13
PART 12 — DIGITAL OPERATIONS / INDUSTRY 4.0	14
PART 13 — HUMAN & BEHAVIORAL OPERATIONS	14
PART 14 — SUSTAINABILITY & SAFETY	14
PART 15 — CORE AXIOMS	15
PART 16 — IMPROVEMENT PHILOSOPHIES	16

Part II — 1,000 Applied Diagnostics (100 SMB Industries) 17

1,000 SMB Diagnostic Questions — Answered Through the Ontology of Operations Management 17

1. Residential remodeling / general contractors	17
2. Plumbing	19
3. HVAC	20
4. Electrical contractors	21
5. Roofing	23
6. Landscaping / lawn care	24
7. Painting contractors	26
8. Pest control	28
9. Janitorial / commercial cleaning	29

10. Handyman services	31
11. Flooring contractors	32
12. Tree services	34
13. Pressure washing	35
14. Water/fire/mold restoration	37
15. Pool service & repair	39
16. Concrete / masonry	40
17. Fencing contractors	42
18. Garage door services	43
19. Appliance repair	45
20. Solar installation	47
21. Dental practices	48
22. Primary care / physician practices	50
23. Chiropractic	51
24. Physical therapy clinics	53
25. Mental health / counseling practices	55
26. Optometry	56
27. Veterinary clinics	58
28. Home health care agencies	59
29. Med spas / aesthetic clinics	61
30. Urgent care clinics	63
31. Law firms (solo/small)	64
32. Accounting / bookkeeping / tax prep	66
33. Marketing / advertising agencies	67
34. IT services / MSPs	69
35. Management consulting	70
36. Staffing / recruiting agencies	72
37. Web design / development shops	73
38. SEO / digital agencies	75
39. HR / payroll services	77
40. Business coaching	78
41. Restaurants (independent)	80
42. Coffee shops & cafes	81
43. Bars & taverns	83
44. Catering companies	85
45. Food trucks	86

46. Bakeries	88
47. QSR franchise units	89
48. Pizza shops	91
49. Ice cream & dessert shops	93
50. Juice / smoothie / boba shops	94
51. Auto repair shops	96
52. Car washes	98
53. Used car dealerships	99
54. Auto detailing	101
55. Tire shops	102
56. Collision / body shops	104
57. Oil change shops	105
58. Towing services	107
59. Transmission shops	109
60. Mobile mechanics	110
61. Hair salons & barbershops	112
62. Nail salons	114
63. Day spas & massage therapy	115
64. Tattoo studios	117
65. Gyms & fitness studios	119
66. Yoga & pilates studios	120
67. Martial arts schools	122
68. Dance studios	123
69. Real estate brokerages	125
70. Property management	126
71. Insurance agencies	128
72. Mortgage brokerages	129
73. Financial advisory firms	131
74. Home inspection	133
75. Title & escrow companies	134
76. Real estate appraisal	136
77. Convenience stores & gas stations	137
78. Clothing boutiques	139
79. Liquor stores	141
80. Hardware stores	142
81. Furniture stores	144

82. Independent pharmacies	145
83. Florists	147
84. Jewelry stores	149
85. Pet supply stores	150
86. Thrift & consignment	152
87. Vape & smoke shops	153
88. E-commerce sellers	155
89. Trucking (owner-operators & small fleets)	157
90. Freight brokerages	158
91. Moving companies	160
92. Courier & last-mile delivery	161
93. Non-emergency medical transport (NEMT)	163
94. Limo & shuttle services	165
95. Childcare & daycare centers	166
96. Tutoring & learning centers	168
97. Self-storage facilities	170
98. Dry cleaners & laundromats	171
99. Private security companies	173
100. Pet grooming & boarding	174
Appendix — Reference Key	176

Part I — The Ontology of Operations Management

The Ontology of Operations Management

PART 0 — FORMAL FOUNDATIONS

0.1 Ontological Commitments

This ontology adopts a **BFO-style upper split** so that time and change are first-class:

Upper Class	Definition	OOM Examples
Continuant	Entity that persists through time	Resource, Material, Facility, Agent, Product, Policy
Occurrent	Entity that unfolds in time; has temporal parts	Process, Activity, Event, State
Event	Point or interval occurrence that changes system state	Breakdown, OrderArrival, Setup, Delivery, Disruption
State	Condition of the system over an interval	Idle, Blocked, Starved, Down, Backordered

Axiom F1 (time-indexing): Every measurable property in OOM is a function of time. $\text{capacity}(r, t)$, $\text{wip}(p, t)$, $\text{demand}(\text{product}, t)$. Any claim made without a time index is shorthand for steady state.

Axiom F2 (dual role): A continuant may be both a **Resource** and a **FlowUnit** — disjointness is defined per role, not per entity. (A patient in a hospital is a flow unit; a surgeon is a resource; a machine under maintenance is temporarily neither-capable.)

0.2 Disjointness & Coverage Axioms

- **D1:** $\text{Process} \subseteq \text{Occurrent}$, $\text{Resource} \subseteq \text{Continuant}$, $\text{Process} \cap \text{Resource} = \perp$
- **D2:** **Input**, **Output**, **Inventory** are roles of **Material** | **Information** | **FlowUnit**, not disjoint classes
- **D3:** **VA**, **NVA**, **NNVA** (value-added status) form a pairwise-disjoint, exhaustive cover of **Activity**
- **D4:** $\text{MTS} \cap \text{MTO} \cap \text{ATO} \cap \text{ETO} = \perp$ at the order-fulfillment-strategy level per product family
- **D5:** Every **Activity** must be assigned ≥ 1 **Resource** (cardinality ≥ 1) and must have ≥ 0 predecessors

0.3 Core Relations (domain $\rightarrow$ range, cardinality)

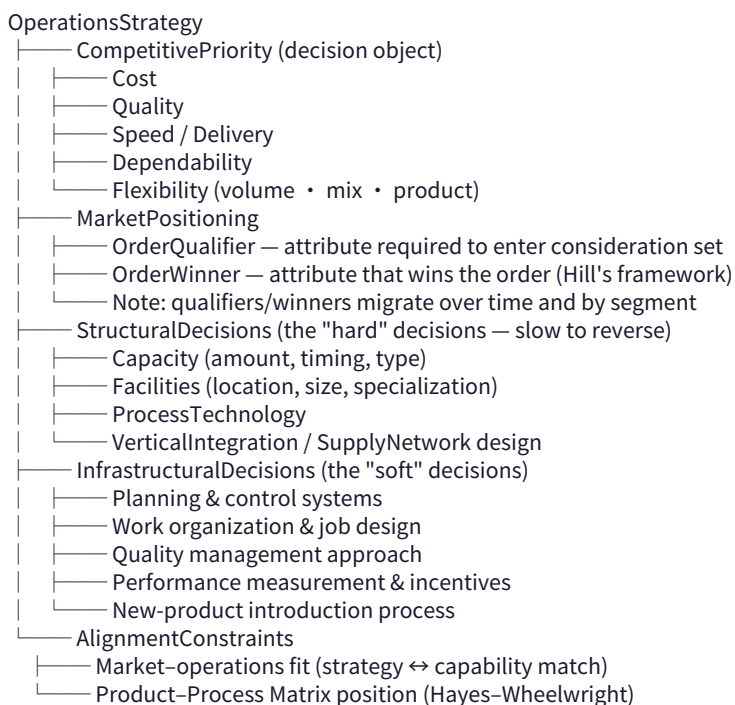
Relation	Domain $\rightarrow$ Range	Card.	Inverse
isPartOf	Activity $\rightarrow$ Process	n:1	hasPart
precedes	Activity $\rightarrow$ Activity	n:n (acyclic)	succeeds
consumes	Activity $\rightarrow$ Resource/Material	n:n	consumedBy

Relation	Domain → Range	Card.	Inverse
produces	Activity → Output/Inventory	n:n	producedBy
transforms	Process → (Input × Output)	n:n	—
assignedTo	Activity → Resource	n:n	executes
constrainedBy	Process → Resource	n:n	constrains
triggers	Event → StateChange	1:n	triggeredBy
fulfills	Process → Demand	n:n	fulfilledBy
specifiedBy	Output → Specification	n:1	specifies
governedBy	Activity/Process → Policy	n:n	governs
measuredBy	Process/Resource/Output → KPI	n:n	measures

PART 1 — STRATEGY LAYER

1.1 Class: OperationsStrategy

The deliberate pattern of decisions that configures operations to support competitive position.



1.2 Key Strategy Classes

- **ProductProcessMatrix:** Job shop → Batch → Line → Continuous mapped against product lifecycle (introduction → growth → maturity). Off-diagonal positions are strategic choices with cost/flexibility consequences.

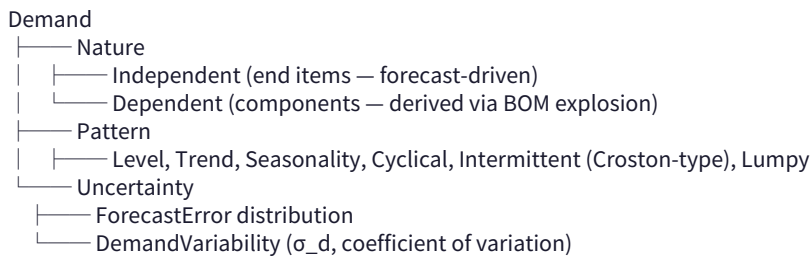
- **Focus:** Plant-within-a-plant; a process serving conflicting priorities underperforms a focused one.
- **CumulativeCapability (Sand Cone) model:** quality → dependability → speed → cost efficiency built cumulatively; contrasts with pure **TradeOff model** (performance frontier shifts only via capability investment).
- **MassCustomization:** strategies — modularity, postponement (form/time/place), ATO architecture, decoupling-point placement.
- **DecouplingPoint:** the inventory position separating forecast-driven (push) from order-driven (pull) stages; position is a strategic variable.
- **LearningCurve / ExperienceCurve:** unit cost or labor hours fall by a constant % per doubling of cumulative volume (typically 70–95% curve slope). $T_n = T_1 \cdot n^{(\log \phi / \log 2)}$.

1.3 Strategy-Level Metrics

- **Productivity:** partial (output / single input), multifactor, total factor
- **BalancedScorecard integration:** financial • customer • internal process • learning & growth
- Operations contribution: **OrderWinner attainment**, **capability maturity**

PART 2 — DEMAND LAYER

2.1 Class: Demand



2.2 Class: ForecastingMethod

Family	Methods	Use
Qualitative	Delphi, market research, sales-force composite, judgmental	New products, long horizon
Time series	Moving average, exponential smoothing, Holt (trend), Holt–Winters (seasonal)	Short/medium horizon, stable patterns
Causal	Regression, econometric, leading indicators	Driver-based demand
ML / demand sensing	Gradient boosting, deep temporal models, POS-signal sensing	High-frequency, data-rich

ForecastError measures: MAD, MSE, MAPE, sMAPE, **bias** and **TrackingSignal** = RSFE / MAD (control limits typically ± 4 to ± 8).

2.3 Demand Management

- **Influencing demand:** pricing, promotions, lead-time quoting, counter-seasonal products
- **Accommodating demand:** backlog, reservation systems, yield-managed allocation
- **Demand shaping input to S&OP** (see Part 7)

PART 3 — PROCESS CORE

3.1 Class: Process (Occurrent)

A structured sequence of activities transforming inputs into outputs of greater value.

Transformation taxonomy:

- **Manufacturing:** fabrication, assembly, continuous processing
- **Service:** front-office (high contact), back-office (low contact), hybrid
- **Logistics:** transport, warehousing, cross-dock, distribution
- **Information transformation:** transactions, claims, software change

Flow structure: Project • Job shop • Batch • Line/Flow • Continuous

Service contact intensity: Pure service • Mixed • Quasi-manufacturing (Chase's customer-contact model: efficiency falls as contact rises)

Order fulfillment strategy: MTS • ATO • MTO • ETO — each implies a decoupling-point position and a lead-time profile.

3.2 Class: Activity (Occurrent)

Atomic unit of work. **Properties:**

Property	Type	Notes
processing_time	stochastic (μ, σ)	variability is first-class
setup_time	stochastic	sequence-dependent setups allowed
predecessors	set	precedence graph must be acyclic
value_added_status	enum	VA / NVA / NNVA — exhaustive (Axiom D3)
assigned_resource	set	cardinality ≥ 1 (Axiom D5)

3.3 Class: FlowUnit (Continuant role)

Entity moving through the process: part, order, batch, customer, patient, transaction, document.

3.4 Process States (Occurrent)

Idle • Blocked (downstream full) • Starved (upstream empty) • Running • Down • Setup — state transitions are **Events** (Part 0.1). Utilization accounting must decompose calendar time across these states.

PART 4 — RESOURCES

4.1 Resource Taxonomy

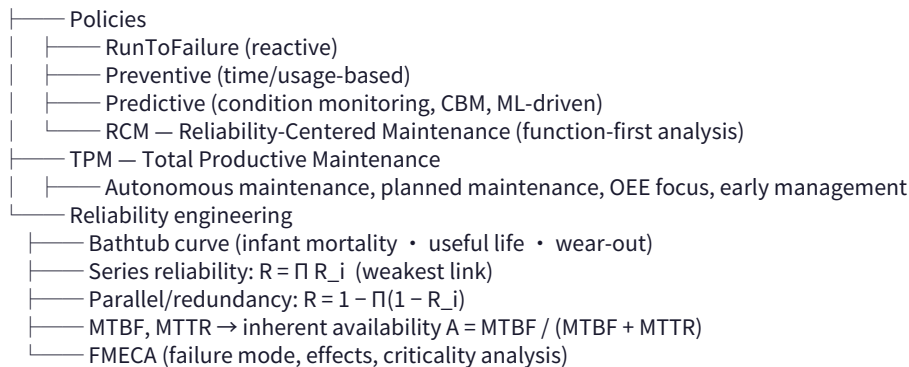
- **Human:** direct labor, indirect labor, supervision, engineering, knowledge work
- **Physical:** machines, tools, fixtures, facilities, automation/robots/cobots, IT/OT systems
- **Materials:** raw, components, consumables, packaging, energy
- **Information:** BOMs, routings, drawings, specifications, recipes
- **Financial:** working capital (inventory + WIP + receivables – payables)

4.2 Capacity — Corrected Formulation

- **Design capacity:** maximum theoretical rate under ideal conditions
- **Effective capacity:** sustainable rate given product mix, schedules, maintenance, quality losses
- **Utilization (corrected):** $U = \text{actual output} / \text{effective capacity}$ — may exceed 1 temporarily against design capacity only via overtime/overclocking, at degradation cost
- **OEE** = Availability × Performance × Quality (six big losses: breakdowns, setups, idling/minor stops, reduced speed, startup rejects, production rejects)

4.3 Class: Maintenance

Maintenance



4.4 Human Resource Properties

- **JobDesign:** specialization ↔ enrichment/rotation; socio-technical systems fit
- **WorkMeasurement:** time study (observed × rating × (1 + allowances)), standard time, predetermined motion-time systems (MTM, MOST), work sampling
- **Ergonomics & safety load:** physical/cognitive demand as capacity constraints
- **Skill matrices:** flexibility of labor as a capacity hedge

PART 5 — INVENTORY

5.1 Inventory Classes (by position & purpose)

- **By position:** RM • WIP • FG • MRO • pipeline/in-transit
- **By purpose:** cycle stock • safety stock • anticipation/seasonal • decoupling • hedge
- **By state:** available • allocated • on-order • backordered • obsolete

5.2 Control Models

Model	Decision rule	When
EOQ	$Q^* = \sqrt{(2DS/h)}$	stable demand, fixed ordering cost
News vendor	$P(\text{Demand} \leq Q^*) = C_u / (C_u + C_o)$	single-period, perishable
(Q, r) continuous review	order Q when position $\leq r$	A items, high service
Periodic review (T, S)	order up to S every T periods	consolidated ordering
Base stock	one-for-one replenishment	low volume, high value
ABC / Pareto classification	manage by value density	portfolio triage

5.3 Multi-Echelon & Risk Pooling

- **Square-root law:** consolidating n independent stocking points into one reduces required safety stock by factor $\sqrt{n}$ (given independent demands)
- **Risk pooling via postponement:** delay differentiation (form/place/time) to pool demand uncertainty
- **Multi-echelon optimization:** echelon stock vs. installation stock policies; service-level allocation across the network

PART 6 — QUALITY MANAGEMENT

6.1 Quality Dimensions (Garvin's 8)

Performance • Features • **Reliability** • Conformance • Durability • Serviceability • Aesthetics • Perceived quality

6.2 Cost of Quality (CoQ)

CoQ

- └── Prevention (training, planning, DFM)
- └── Appraisal (inspection, testing, audits)
- └── Internal failure (scrap, rework, downtime)
- └── External failure (warranty, returns, recalls, reputation)

Classical result: prevention/appraisal investment shifts the total-cost optimum toward lower defect rates; Six Sigma pushes the curves themselves down.

6.3 Quality Methods

- **SPC:** control charts (X-R, p, c, u, EWMA, CUSUM), common vs. special cause variation, process capability $C_p = (USL - LSL) / 6\sigma$, $C_{pk} = \min(C_{pu}, C_{pl})$ accounting for centering
- **Acceptance sampling:** OC curves, AQL, LTPD, single/double/sequential plans
- **Poka-yoke:** error-proofing at source (contact / fixed-value / motion-step methods)
- **QFD:** House of Quality translating voice-of-customer → engineering characteristics
- **Standards:** ISO 9001 (QMS), ISO 14001 (environment), industry regimes (IATF 16949, AS9100, GMP)
- **Service quality:** SERVQUAL gaps model (knowledge • standards • delivery • communication • service gaps); dimensions: tangibles, reliability, responsiveness, assurance, empathy

PART 7 — PLANNING & CONTROL HIERARCHY

Horizon	Class	Content	Key methods
Strategic (1–5 yr)	Capacity & network design	facility location, sizing, technology, vertical integration	factor rating, center of gravity, transportation/opimization models, real options
Tactical (3–18 mo)	S&OP / aggregate planning	demand–supply balance at family level	chase vs. level vs. hybrid, overtime/subcontract/backlog levers
Medium (weeks)	MPS / MRP	end-item schedule → BOM explosion → net requirements	lot sizing (L4L, EOQ, POQ), safety stock, safety lead time
Short (days)	Scheduling & sequencing	resource-level task ordering	dispatch rules (SPT, EDD, CR, ATC), finite-capacity scheduling, TOC drum-buffer-rope, shifting-bottleneck heuristics
Real-time	Execution & control	release, monitor, intervene	Kanban/CONWIP pull signals, MES, Andon, SPC alarms

Project scheduling (Part 9) plugs in here for project-type processes.

PART 8 — SERVICE OPERATIONS

8.1 Queuing Theory (formal class)

- **Kendall notation** A/S/c/K/N/D (arrival / service / servers / capacity / population / discipline)
- **Canonical results:** M/M/1 → $L = \lambda / (\mu - \lambda)$, $W = 1 / (\mu - \lambda)$; M/M/c → Erlang C (probability of waiting) — the staffing engine for call centers

- **Kingman (VUT) equation:** $CT_q \approx ((c_a^2 + c_e^2)/2) \cdot (U^{\sqrt{2(m+1)}-1}/(m(1-U))) \cdot t_e$
— waiting time explodes with utilization **and** variability; this is the quantitative core of the variability-buffering law

8.2 Service-Specific Structures

- **Service blueprinting:** line of interaction / visibility; fail points mapped to poka-yokes
- **Service recovery:** recovery paradox, complaint handling as process
- **Revenue / yield management:** pricing perishable capacity; fences, overbooking, littlewood's rule, EMSR-b seat/protection levels — applies to airlines, hotels, rental, advertising inventory, healthcare slots

PART 9 — PROJECT MANAGEMENT

ProjectManagement

- Network methods: CPM (deterministic), PERT (β -distributed times, expected path)
- Critical path & slack (total, free)
- Time-cost trade-off: crashing (cost slope per activity)
- Critical Chain (Goldratt): feeding buffers, project buffer, resource contention
- Earned Value Management: PV, EV, AC $\rightarrow$ SPI = EV/PV, CPI = EV/AC, EAC
- Risk: PERT Monte Carlo, schedule risk analysis

PART 10 — FACILITY & LAYOUT DESIGN

- **Layout types:** fixed-position • process/functional • product/line • cellular (GT-based) • hybrid
- **Line balancing:** given takt time and precedence graph, minimize stations / balance delay;
theoretical min stations = $\sum t_i / \text{takt}$
- **Cellular design:** part-family formation, machine cells, virtual cells
- **Location models:** factor rating, center of gravity, transportation LP, coverage models for services

PART 11 — SUPPLY CHAIN & SOURCING

11.1 SCOR Processes

Plan • Source • Make • Deliver • Return • **Enable**

11.2 Sourcing & Procurement

- **Kraljic matrix:** leverage / strategic / bottleneck / non-critical items — by profit impact $\times$ supply risk
- **Sourcing structure:** single vs. dual vs. multi-source; global vs. local; make-vs-buy (transaction cost + capability analysis)
- **Supplier management:** segmentation, development programs, scorecards, contracts (revenue-sharing, buy-back, quantity flexibility — coordination vs. double marginalization)

11.3 Systemic Phenomena

- **Bullwhip effect:** order variance amplifies upstream; causes — demand signal processing, rationing game, order batching, price fluctuation; countermeasures — information sharing (CPFR), everyday-low-pricing, lead-time reduction, vendor-managed inventory
- **Supply chain coordination failure:** double marginalization, incentive misalignment

11.4 Risk & Resilience

RiskManagement

- Disruption classes: supply, demand, logistics, facility, geopolitical, cyber, pandemic
- Assessment: probability × impact mapping, value-at-risk, time-to-recover (TTR) vs. time-to-survive (TTS)
- Mitigation: redundancy, flexibility, buffering, visibility (control towers), supplier diversification, nearshoring
- Business continuity planning (BCP)
- Resilience metrics: recovery time, robustness vs. efficiency frontier

PART 12 — DIGITAL OPERATIONS / INDUSTRY 4.0

DigitalOps

- Sensing layer: IoT, RFID, machine telemetry, computer vision QC
- Platform layer: MES, ERP, APS, WMS, control towers
- Model layer: digital twins (asset • process • network scale)
- Intelligence layer:
 - Predictive maintenance (ML on condition data)
 - Demand sensing & ML forecasting
 - RL/heuristic scheduling
 - Anomaly detection for quality
- Automation: fixed • flexible • collaborative robots (cobots) • AGV/AMR
- Additive manufacturing:
 - Strategic effect: collapses setup cost → breaks product–process matrix (variety without batching penalty); enables distributed/spare-part-on-demand networks

PART 13 — HUMAN & BEHAVIORAL OPERATIONS

- **Behavioral ops:** how judgment deviates from model assumptions — newsvendor pull-to-center bias, bullwhip overreaction, queue abandonment psychology (Balkema), incentive-induced gaming of KPIs (Goodhart's law)
- **Work organization:** teams, cell ownership, empowerment as jidoka prerequisite
- **Incentives & metrics design:** individual vs. group, gainsharing, avoiding local-optima suboptimization
- **Learning & knowledge:** learning curves (Part 1.2), standard work as codified knowledge, A3 thinking, communities of practice

PART 14 — SUSTAINABILITY & SAFETY

- **Triple bottom line:** people • planet • profit as co-objectives
- **Green operations:** carbon/water/waste footprint per unit, energy-efficient scheduling, closed-loop / circular supply chains, remanufacturing, industrial symbiosis

- **Compliance & standards:** ISO 14001, ISO 45001 (safety), ESG reporting integration into ops KPIs
- **Safety management:** hazard analysis, OSHA regimes, safety stock of attention — leading indicators (near-misses) vs. lagging (incident rates)

PART 15 — CORE AXIOMS

A1 — Little's Law (with preconditions):

$WIP = \lambda \cdot W$ for any subsystem in **stable equilibrium** (long-run flow balance: $\lambda_{in} = \lambda_{out}$, no perpetually growing queue). Holds regardless of arrival/service distributions.

A2 — Throughput bound (corrected):

$Throughput \leq \min(BottleneckCapacity, Demand)$

A demand-constrained process cannot be diagnosed with capacity logic alone.

A3 — Bottleneck-set formulation:

$B(t) = \operatorname{argmin}_{\{r \in Resources\}} effective_capacity(r, t)$

- B may be a **set** (ties, parallel identical resources)
- B is **time-varying**: product-mix changes shift the argmin (shifting bottleneck)
- Invariant: $Throughput(p, t) \leq capacity(b, t)$ for any $b \in B(t)$

A4 — Variability-buffer law (qualitative Kingman):

Waiting time increases non-linearly as utilization $\rightarrow 1$ and as arrival/process variability (c_a^2 , c_e^2) increase. A system with variability requires buffer(s): **inventory, capacity, or time** — choose deliberately; one is always paid.

A5 — Conservation of flow:

$inflow(t) - outflow(t) = d(inventory)/dt$ — no exceptions; backorders relax the inventory ≥ 0 constraint at penalty cost.

A6 — OEE decomposition:

$EffectiveOutput = CalendarTime \times Availability \times Performance \times Quality$

A7 — Risk pooling (square-root law):

Consolidating n i.i.d. demand streams: $\sigma_{pooled} = \sigma/\sqrt{n}$ -equivalent safety stock.

A8 — Traceability:

Every output must trace to ≥ 1 input via an activity chain (**producedBy** • **consumes** closure is non-empty).

A9 — Learning curve:

$T_n = T_1 \cdot n^{-b}$, $b = \log(\phi)/\log(2)$, $\phi \in (0,1]$ — cumulative volume reduces unit time predictably in labor-intensive processes.

A10 — Strategy-structure consistency:

For every [OrderWinner](#) there must exist ≥ 1 structural/infrastructural decision configured to deliver it (orphaned strategy = invalid state).

PART 16 — IMPROVEMENT PHILOSOPHIES

Philosophy	Core mechanism	Links to
Lean / TPS	eliminate muda (7+1 wastes), JIT, jidoka, heijunka, standard work, kaizen	Parts 3, 4.3, 13
Six Sigma	variation reduction, DMAIC / DMADV	Parts 6, 15-A4
TOC	5 focusing steps; exploit → subordinate → elevate; drum-buffer-rope	Parts 7, 15-A3
TQM	PDCA, QFD, company-wide quality culture	Part 6
SCM / SCOR	Plan-Source-Make-Deliver-Return-Enable	Part 11
Resilience engineering	anticipate • monitor • respond • learn	Part 11.4

Part II — 1,000 Applied Diagnostics: 100 SMB Industries × 10 Questions

1,000 SMB Diagnostic Questions — Answered Through the Ontology of Operations Management

How to read this book: each answer first names the operational mechanism the question is circling, then the directive that follows from the ontology in Part I.

Citation key: every answer cites its ontology anchors in brackets — § n = Part/Section of Part I, A1–A10 = Axioms (Part 15), D1–D5 = Disjointness Axioms, F1–F2 = Formal Axioms.

1. Residential remodeling / general contractors

Q1. Change-order disputes cluster on jobs where clients selected finishes before demolition revealed structural conditions. Contract language or discovery process?

A: The discovery process. Finish selection is an [Activity](#) scheduled before its true predecessor — demolition-based discovery — could define its input [Specification](#); the precedence graph is broken, so the contract is asked to absorb what sequencing should have prevented. Make demolition a hard predecessor with a priced allowance gate. [§ 3.2 predecessors, § 9 CPM, D5, A8]

Q2. When a project goes sideways, is the first signal in job-costing reports or the client's tone of voice?

A: The client's tone. Job-costing reports are lagging [measuredBy](#) KPIs on completed work; the client's tone is a real-time [Event](#) signaling a [StateChange](#). The fix is not to abandon cost reports but to instrument leading process states (schedule slip, RFIs aging) so control is not outsourced to the customer's emotions. [§ 0.3 measuredBy, § 0.1 Event/State, § 7 real-time execution & control]

Q3. Referrals from past clients convert far better than realtor referrals, yet the relationship budget goes to realtors. What is the budget optimizing for?

A: Measurability, not conversion — a textbook Goodhart distortion: the channel with an attributable line item wins the budget regardless of order-winning power. Reallocate by conversion-weighted margin per channel, not spend visibility. [§ 13 incentive-induced gaming / Goodhart, § 1.1 OrderWinner, § 1.3 performance measurement]

Q4. Our best carpenters produce their worst work in the final 10% of jobs. What does the punch-list sequence reward?

A: Starting, not finishing. Payment and scheduling milestones treat substantial completion as the terminal [Event](#), so the punch list is [NVA](#) work nobody is incentivized to close; earned value is declared before value is actually earned. Tie final payment and crew release to punch-list closure. [D3 VA/NVA, § 13 incentives, § 9 EVM]

Q5. Which subs cause cascading schedule slips we absorb as reputation damage instead of billing back, and why have we never measured it?

A: Because no supplier scorecard exists — slip propagation is an unmeasured [external failure](#) cost transferred to you by a coordination failure. Assign each sub a delay-cost ledger; what gets [measuredBy](#) gets negotiated. [§ 11.2 supplier management/scorecards, § 6.2 CoQ external failure, A8 traceability]

Q6. Would margins improve more from an 8% price rise or firing bottom-decile clients, and why can't our books answer that?

A: The books lack customer-level cost allocation — revenue per client is visible, consumed [Resource](#) hours per client are not, so client-level margin is uncomputable. Implement job-costing by customer before choosing; typically firing the bottom decile wins because it also frees bottleneck crew capacity. [§ 4.1 Financial resource, § 1.3 productivity metrics, § 0.3 measuredBy]

Q7. Close rate collapses when the homeowner's spouse is not at the estimate. Who is the proposal written to convince?

A: Only half the buying unit. The proposal is a one-sided [Specification](#) for a two-agent decision; the absent spouse later re-opens it, which reads as a stalled close. This is the SERVQUAL knowledge gap — you misread who the customer is. Route proposals through both decision-makers or design them to be re-presentable. [§ 6.3 SERVQUAL knowledge gap, § 2.3 demand management]

Q8. The production calendar jams every Sep–Oct while the pipeline looks healthy year-round. Where does the handoff lose the timing?

A: At the missing tactical layer: sales promises start dates against infinite capacity because no S&OP-style aggregate plan converts pipeline into finite crew-weeks. Seasonal demand is a known [Pattern](#); the jam is a capacity buffer never deliberately chosen. [§ 7 S&OP/aggregate planning, § 2.1 Demand Pattern seasonality, A4]

Q9. Where does warranty call-back money leak: workmanship, or scope documents that never defined done?

A: Predominantly the scope documents. "Done" undefined means the [Output](#) has no enforceable [Specification](#), so every call-back is arguable; workmanship defects would cluster by crew, scope defects cluster by contract language. Trace claims via [producedBy](#) to their spec source before retraining anyone. [§ 0.3 specifiedBy, § 6.2 external failure, A8]

Q10. Our estimating model assumes margins grow with project size; past \$150k they fall. Where does the model stop being true?

A: At the point where coordination complexity outruns the learning curve — larger jobs move you off your Product–Process Matrix sweet spot into longer critical paths, more interfaces, and more rework loops. The model fails where project risk stops scaling linearly with revenue. Re-price using size-conditional historical margins. [§ 1.2 ProductProcessMatrix, A9 learning curve limits, § 9 schedule risk]

2. Plumbing

Q1. We win some neighborhoods and lose adjacent ones at identical pricing. What differs that pricing cannot see?

A: Order winners differ by micro-segment — response time reputation, truck visibility, one anchor customer's advocacy. Price is a qualifier; the winner varies block by block with referral-network topology. Map wins/losses against referral sources per neighborhood. [§ 1.1 OrderQualifier vs OrderWinner, § 2.1 Demand Uncertainty]

Q2. Our review score drops exactly when monthly revenue peaks. Capacity strain or a CSR handling pattern?

A: Test capacity first: at peak revenue, utilization approaches 1 and Kingman's law makes wait and error rates explode non-linearly — the review drop is the variability-buffer law presenting as a service problem. If CSR quality only degrades under load, it is still capacity. Measure wait-time and callback variance against utilization. [§ 8.1 Kingman VUT, A4, § 6.3 SERVQUAL reliability]

Q3. Drain-cleaning calls feed almost no upsells; competitors' do. Diagnostic skill or incentive design?

A: Incentive design until proven otherwise: techs paid on tickets-per-day treat diagnostics as [NVA](#) delay. Skill is trainable only after the scorecard stops punishing the diagnostic pause. Change what [measuredBy](#) rewards, then audit skill. [§ 13 incentives & Goodhart, D3, § 0.3 measuredBy]

Q4. Why do commercial property managers churn at ~18 months with zero documented complaints?

A: Zero complaints means silence, not satisfaction — the knowledge gap: you never measured their decision criteria, and the renewal was never managed as an [Activity](#) with an owner. Build a renewal process that starts at month 12. [§ 6.3 SERVQUAL knowledge gap, D5 every activity needs a resource]

Q5. Which service-agreement benefits drive renewals versus which we advertise because competitors do?

A: Run QFD in reverse: correlate renewal against benefit usage data, not benefit presence in the brochure. Benefits nobody uses are qualifiers you've mistaken for winners. [§ 6.3 QFD voice-of-customer, § 1.1 OrderQualifier]

Q6. When parts costs spike, do we reprice within days or absorb until margin alarms fire?

A: You absorb — the alarm is a lagging control. Supplier price fluctuation is a known bullwhip cause; the countermeasure is an indexed price-book with automatic pass-through triggers, i.e., a [Policy](#) governing the repricing [Activity](#). [§ 11.3 price fluctuation, § 0.3 governedBy, F1 time-indexing]

Q7. Callbacks cluster on two technicians: skill, parts, or dispatch expectations?

A: Treat it as SPC: clustering on two techs out of many is special-cause variation by definition. Decompose before acting — callback coded by failure mode will separate workmanship (skill) from parts brand from promise-setting (dispatch). [§ 6.3 SPC common vs special cause, A8 traceability]

Q8. Does flat-rate pricing lose complex jobs to hourly shops on perceived honesty?

A: Yes, in the segment that reads flat-rate as a black box. Perceived quality is a Garvin dimension independent of actual quality; for complex jobs, itemized transparency is the order winner. Offer both frames and let segment data decide. [§ 6.1 perceived quality, § 1.1 OrderWinner, § 2.3]

Q9. Emergency calls should carry urgency pricing, yet their average ticket runs below scheduled work. What discounts urgency at the point of sale?

A: The CSR/tech, who gives the price before the customer reveals the urgency premium they'd pay — you have perishable, time-sensitive capacity and no yield management. Price the response-time tier explicitly; urgency is a product attribute, not a favor. [§ 8.2 revenue/yield management, § 2.3 influencing demand]

Q10. At what call volume does adding a truck stop paying, and why does our math ignore dispatch?

A: Because the model assumes trucks are the bottleneck when past ~4–5 trucks the constraint migrates to dispatch coordination. The bottleneck is a time-varying set; elevate dispatch (better scheduling, zones) before buying truck N+1. [A3 bottleneck set, § 7 finite-capacity scheduling, § 16 TOC]

3. HVAC

Q1. Why do maintenance-agreement customers generate less annual revenue per home than non-agreement customers?

A: Self-selection plus pricing: agreements attract the price-sensitive and are priced as a discount bundle rather than a relationship that funds priority service and replacement pipeline. Decompose revenue into service vs. replacement; if agreement homes also replace less, the product is subsidizing the wrong behavior. [§ 2.3 demand shaping, § 1.1 CompetitivePriority, § 13]

Q2. When the first heat wave hits, which fails first: phones, techs, or warehouse?

A: Whichever has the lowest series reliability — a surge exposes the weakest link in a serial system, and it is usually phones (unbuffered arrival spike) before techs. Compute each link's capacity against the surge rate and pre-position buffers deliberately: inventory, capacity, or time — you will pay one. [§ 4.3 series reliability, A4 variability–buffer law, § 8.1 Erlang C for phones]

Q3. More summer no-cool calls from equipment age or our install quality 5–7 years ago?

A: Both live on the bathtub curve, but only one is yours: install-quality failures surface at exactly 5–7 years as early wear-out, and they are traceable. Tag no-cool calls by install crew and vintage; if your work is the cause, that's external failure cost coming home. [§ 4.3 bathtub curve, § 6.2 external failure, A8]

Q4. Does seasonal installer hiring create a Q3 callback dip?

A: Almost certainly: new hires sit at the top of the learning curve, raising process variability precisely when volume peaks — a compounding Kingman effect. Track callback rate by installer tenure; if confirmed, the fix is skill-matrix staffing, not blanket retraining. [A9 learning curve, § 4.4 skill matrices, A4]

Q5. Why does replacement close rate drop when the comfort advisor has <6 months tenure — skill or undocumented quoting process?

A: Undocumented process. If close rate depended on teachable skill, training would fix it; depending on tenure means the quoting logic lives in veterans' heads — standard work never codified. Codify the quote, then tenure stops mattering. [§ 13 standard work as codified knowledge, § 16 Lean standard work, A9]

Q6. We lose replacement bids to competitors whose equipment costs more. Is it financing presentation?

A: Yes. The order winner for a \$8–12k replacement is the monthly payment and approval certainty, not the box price. You're competing on qualifier while they win on the winner. Lead with the financed monthly number. [§ 1.1 OrderWinner migration, § 2.3 influencing demand]

Q7. Are agreements priced to retain customers or to subsidize replacements — and which does renewal data show?

A: Let the renewal cohort answer: if agreement customers replace at higher rates, it's a working loss-leader (keep it); if they renew but never replace, it's a discount club (reprice). The question is empirical; the data exists in the renewal-to-replacement conversion funnel. [§ 2.1 Demand pattern analysis, § 1.3 metrics, § 13]

Q8. Why does commercial PM renew at 90% while residential renews at 55%?

A: Different buyers and different switching costs: commercial buys risk transfer with a budget line (high dependability value), residential buys a discount they re-evaluate emotionally each year. Same product, two order winners — price and sell them as two products. [§ 1.1 OrderWinner by segment, § 6.3 SERVQUAL assurance]

Q9. Why do certain zip codes produce 2.5× the warranty claims per install?

A: Special-cause variation waiting to be decomposed: installer crew assignment, housing stock age, water quality, or permit inspection rigor all correlate with geography. Stratify claims by crew × zip; if crew disappears as a factor, it's the stock. [§ 6.3 SPC stratification, A8, F1]

Q10. Tens of thousands in zero-movement parts while we place weekly emergency orders. Who owns the stocking logic?

A: Nobody — the stocking logic is an orphaned [Activity](#) with no assigned [Resource](#). The fix is an ABC/criticality classification: A-items by downtime cost (not unit cost) get stocking rules; the tail gets zero-stock with emergency freight priced into the flat-rate book. [D5 cardinality, § 5.2 ABC classification, § 5.1 MRO, § 0.3 measuredBy]

4. Electrical contractors

Q1. Why does emergency revenue spike in specific subdivisions — deliberate marketing or reacting to infrastructure age?

A: You're reacting. Subdivision clusters map to construction vintage — aluminum branch wiring and original panels hitting the wear-out region together. That's a causal forecasting gift: map permits by year, market to the bathtub curve before the failure, and the emergency premium becomes planned replacement work. [§ 4.3 bathtub curve, § 2.2 causal forecasting/leading indicators]

Q2. Do apprentices turn net-positive by month 18, or are we a training program for whoever poaches them at month 24?

A: Compute it honestly: apprentice cost vs. billed contribution is negative until mid-curve; if attrition hits at month 24, you capture the investment and competitors capture the return. Either retention economics (pay at curve inflection) or faster billing utilization must change — otherwise you are, in fact, a school. [A9 learning curve, § 4.4 skill matrices, § 13]

Q3. Journeyman utilization drops 20% when we run more than four trucks. Dispatch design or job mix?

A: Dispatch design. Coordination load grows combinatorially with trucks while job mix doesn't change discontinuously at truck five; the constraint migrated from field capacity to the dispatcher. Add scheduling structure (zones, finite-capacity rules) before concluding anything about mix. [A3 shifting bottleneck, § 7 dispatch rules/finite scheduling]

Q4. When commercial tenants churn, how much tenant-improvement revenue evaporates — and do we see it coming?

A: You don't, because TI demand is dependent demand derived from your GCs' and property managers' pipelines, which you don't instrument. A leading indicator (tenant lease expirations in served buildings) converts surprise churn into forecast. [§ 2.1 dependent demand, § 11.4 disruption assessment, § 2.2 leading indicators]

Q5. GCs keep us on bid lists but stop awarding after three consecutive wins. What are we signaling?

A: Probably capacity saturation — the third win signals you'll be stretched on the fourth, so they protect their schedule. Either your backlog is visible and unmanaged, or win behavior changes your bid sharpness. Track bid margin vs. backlog to see which. [A3 bottleneck, § 13 behavioral ops, § 1.1 OrderQualifier dependability]

Q6. Why does quote acceptance fall when we itemize versus bundle — and which segment drives it?

A: Itemizing lets residential customers price-shop line items and anchors them on the visible ones; commercial buyers often require itemization. Segment it: bundle for homeowners (sell the outcome), itemize for GCs (sell the scope). One format cannot win both order winners. [§ 1.1 OrderWinner by segment, § 6.1 perceived quality]

Q7. Our two best estimators disagree by 15% on identical job walks. Whose takeoff is the outlier and what's the annual cost?

A: Neither is "the outlier" — you have no takeoff standard, so you have measurement-system variation, not estimator error. Build the standard work (unit prices, waste factors, labor units per device), then measure both against it; the 15% spread is your annual cost until then. [§ 13 standard work, § 4.4 work measurement, § 6.3 SPC]

Q8. Is code-update work (GFCI/AFCI) a systematically harvested lead source or left to inspectors to trigger?

A: Left to inspectors. Code changes are scheduled, published demand events — a deterministic demand signal you can plan campaigns around. Treat each code cycle as a product launch with its own mini S&OP. [§ 2.3 demand management, § 2.1 Demand Nature, § 7 S&OP]

Q9. Residential margin erosion: copper pass-through failure or un-requested scope creep?

A: Decompose the ledger: copper volatility shows in material-cost variance; scope creep shows in labor-hour overrun on fixed quotes. Both are policy failures — no indexed material pass-through, no change-order trigger at the service-call level. Measure each for one quarter; the answer splits cleanly. [§ 11.3 price fluctuation, § 9 EVM-style variance, § 0.3 governedBy]

Q10. Panel-upgrade leads from EV installs convert at half the standalone rate. What does the EV customer already believe?

A: That the charger works without the upgrade — the sale already "succeeded," so the panel reads as an upsell, not a prerequisite. The standalone inquirer already accepts the panel as the product. Reframe at quote time: the deliverable is charging capacity, and the panel is part of the [Specification](#), not an option. [§ 6.3 SERVQUAL knowledge gap, § 0.3 specifiedBy, § 1.1 OrderQualifier]

5. Roofing

Q1. Why do we win hail-claim work in adjacent counties but lose it in our home county?

A: Home-county order winners differ — likely adjuster relationships and incumbent storm-chaser saturation versus your local retail reputation, which doesn't transfer into insurance work. Your brand wins replacements; it doesn't win claims. Build the adjuster/agent channel deliberately or accept the split. [§ 1.1 OrderWinner by segment, § 11.2 relationship structure]

Q2. Why don't neighborhoods we roofed 8–10 years ago generate the replacement wave the installed base predicts?

A: Because the installed base is a latent demand asset you never instrumented: no recontact cycle, so when the wear-out window opens, the customer calls whoever is visible that month. The installed base is your best leading indicator — work it like a crop with a harvest calendar. [§ 2.2 leading indicators, § 4.3 bathtub/wear-out, § 2.3]

Q3. Commercial flat-roof book growing while residential shrinks — strategy or accident?

A: Accident until someone can state the intended mix. A drift in mix is an emergent strategy; the ontology requires every order winner to be backed by a deliberate structural decision, and here none was made. Decide: the two books need different crews, sales motion, and cash-cycle tolerance. [§ 1.1 OperationsStrategy, A10 strategy–structure consistency]

Q4. Which kills more margin: weather days or material deliveries arriving after the crew?

A: The deliveries. Weather is exogenous variability you buffer with schedule slack; a crew waiting on a truck is a self-inflicted starved state with full labor cost burning — a supplier-coordination failure you can actually fix with delivery-window policies and staging. [§ 3.4 Starved state, § 11.2 supplier management, A4]

Q5. What makes a cash buyer walk: price, or a proposal written in insurance language?

A: The proposal. Insurance language signals "this isn't for you" — a communication gap in SERVQUAL terms. The cash buyer needs scope, materials, and warranty in consumer terms; write two templates and route by payer type. [§ 6.3 SERVQUAL communication gap, § 6.1 perceived quality]

Q6. Why do storm-season insurance jobs close fast but produce our worst margins and worst reviews simultaneously?

A: Surge economics: volume spikes, utilization hits 1, and variability law takes over — crews rush, cycle quality drops, supplements get sloppy, and you discount to keep pace with storm-chasers. You're running emergency capacity at everyday prices. Storm work needs its own price book and capacity buffer. [A4, § 8.2 yield management, § 6.2 external failure]

Q7. When we sub out storm overflow, what does the subbed-work callback rate cost in reputation versus the revenue it saves?

A: You don't know, which is the problem: subcontracted work carries your brand with their process control — external failure risk outsourced but reputation retained. Put subs under the same callback measurement as crews and price the overflow work to include that risk premium. [§ 6.2 external failure, § 11.2 supplier scorecards, A8]

Q8. Supplement recovery varies 40% between estimators on identical claims. What do high-recovery estimators document differently?

A: Photo density, line-item specificity, and code citations — i.e., they write to the adjuster's audit checklist. That's codifiable standard work: extract the high performers' documentation template, make it the spec, and the 40% spread becomes a training exercise. [§ 13 standard work/A3 thinking, § 0.3 specifiedBy, § 4.4 work measurement]

Q9. Crew productivity drops on architectural shingle jobs versus what labor rates assume. Which labor table needs rewriting?

A: The one that treats "shingle" as a single work standard. Architectural shingles change handling, cutting, and fastening times — your standard time data came from 3-tab era jobs. Re-run work measurement by product type; estimating inherits whatever the labor table believes. [§ 4.4 work measurement/standard time, § 13 standard work]

Q10. Does warranty registration predict referral velocity — and why is registration below 30%?

A: Registration is a post-sale touchpoint that reopens the relationship; customers who complete it are self-selected advocates, so yes, it likely predicts referrals. It's below 30% because it's designed as paperwork (manufacturer's process) instead of a relationship event with an incentive attached. Redesign it as the first step of the referral program. [§ 2.3 demand management, § 8.2 service blueprinting, § 13]

6. Landscaping / lawn care

Q1. Does our snow-removal book subsidize spring startup or cannibalize its capacity every March?

A: Both are computable — allocate March crew-hours and equipment between snow standby and spring mobilization. Snow is a counter-seasonal demand-shaping product (good strategy), but if March call-backs collide with plowing standby, you never resolved the capacity conflict on paper. Chase vs. level: pick which book gets the crews in the overlap weeks. [§ 2.3 counter-seasonal demand, § 7 chase vs. level, A3]

Q2. Equipment downtime clusters at the same two points every season despite a written maintenance schedule. What is it missing?

A: Usage-based triggers. A calendar schedule ignores that those two points follow peak-load periods (spring cleanup, fall leaf season); failures are wear-out after stress, not after dates. Switch to condition/usage-based PM keyed to engine hours and post-peak inspections. [§ 4.3 Preventive vs Predictive maintenance, § 4.3 bathtub curve wear-out]

Q3. Which crew configuration maximizes revenue per labor hour at our route density — and have we tested an alternative?

A: You haven't — configuration is an untested [StructuralDecision](#). Run a designed experiment: 2-man vs. 3-man crews on matched routes for four weeks, measuring revenue per crew-hour (not per day). Density changes the answer, so test at your actual density, not the industry's. [§ 1.1 StructuralDecisions, § 1.3 partial productivity, § 4.4 work measurement]

Q4. When clients compare our design-build quotes to pool and deck builders, what must our pricing answer?

A: The capital-project question: financing, phasing, and total project cost — because the customer re-segmented you into a different consideration set with different order winners. Either sell like a capital project (drawings, financing, phased scope) or redirect the comparison back to landscape outcomes. [§ 1.1 OrderQualifier/consideration set, § 2.3]

Q5. Spring maintenance contracts cancel at twice the rate of fall signings. What does the spring customer buy that the fall customer does not?

A: An impulse cure for winter guilt — emotional demand that decays with the season. The fall signer bought planning. Spring signings need an engineered commitment device (auto-renew, preseason perks) because the demand itself is more perishable than the service. [§ 2.1 Demand Uncertainty, § 13 behavioral ops, § 8.2]

Q6. Our best project leads come from properties we do not maintain. What do our clients see that strangers do not?

A: Your maintenance work — which apparently doesn't advertise design capability. Your crews produce invisible quality (reliability), while prospects buy visible transformation. The maintenance base is an under-exploited demand asset: signage, portfolio leave-behinds, and crew-driven enhancement flags convert it. [§ 1.1 MarketPositioning, § 2.3, § 6.1 aesthetics vs reliability]

Q7. Why does enhancement revenue per maintenance client stall after year two — saturation, or do we stop asking?

A: Test the asking first: plot enhancement offers made per client by tenure. In almost every route business, offer frequency decays with familiarity long before property potential does. Saturation is a property-level fact; silence is a process-level habit. [§ 2.3 influencing demand, § 13 incentives, § 0.3 measuredBy]

Q8. Does route density or crew skill explain more variance in per-property margin — and why have we never split them?

A: Because they arrive entangled: dense routes go to experienced crews. Split with a variance decomposition — margin $\sim$ density + crew + interaction over one season of route data. Typically density dominates drive-time variance while skill dominates quality callbacks; the split tells you whether to buy routing software or training. [§ 10 location/routing models, § 6.3 SPC variance decomposition, F1]

Q9. Why hold pricing flat for three-year clients while fuel and labor compound — is the churn risk real or assumed?

A: Assumed. You've never measured price elasticity of your tenured base, so an untested belief governs pricing policy. Run a controlled increase on a cohort; churn response to a 4–6% annual escalator in contract services is usually small, and compounding cost absorption is certainly fatal. [§ 2.3 pricing as demand lever, § 0.3 governedBy Policy, § 13]

Q10. When a property manager switches us out, what did our last 90 days of service data look like — were we coasting?

A: Usually yes: incumbency breeds a slow service decay invisible to you because no complaint fired. Track leading indicators per account (missed-detail rate, response latency, photo documentation) — churn at renewal is decided a quarter earlier, in your own log. [§ 6.3 SERVQUAL reliability, § 7 real-time control, § 12 sensing]

7. Painting contractors

Q1. Does warranty claim rate correlate with paint line, prep checklist, or crew lead?

A: Stratify the claims by all three — that's basic SPC factor separation. In coatings, prep (process adherence) and crew lead (special cause) dominate materials in claim variance; the paint line is usually the smallest factor. Verify, don't assume. [§ 6.3 SPC stratification, § 6.2 external failure, A8]

Q2. Why does estimate accuracy degrade on pre-1970 homes — what prep variable are we underpricing?

A: Substrate uncertainty: lead paint, plaster condition, and layered failing coatings turn prep from a known task into a discovery process. Your estimate prices the mean; pre-1970 prep has a fat right tail. Add a substrate-discovery contingency line or price the tail explicitly. [§ 2.1 Demand/Process Uncertainty, § 9 risk analysis, § 4.4 work measurement]

Q3. When we lose bids, is it the number or the five-day turnaround while competitors quote on-site?

A: Run the loss log by quote lag: in residential services, speed is an order winner independent of price — a same-day quote signals the dependability the customer is actually buying. If losses concentrate above 48 hours, the fix is estimating capacity, not price. [§ 1.1 OrderWinner speed, § 8.1 waiting psychology, § 2.3]

Q4. Do color-consultation add-ons increase close rate enough to offset the extra trip?

A: Compute conversion lift $\times$ margin vs. trip cost. The consultation is a demand-shaping investment: if it moves close rate materially (it usually does on exteriors and cabinets), it's not a cost center, it's the selling activity itself. If it doesn't, you're performing free design. [§ 2.3 demand shaping, § 1.3 productivity, D3]

Q5. Interior closes at 60%, exterior at 30% from identical lead sources. What does the exterior estimate fail to do?

A: Overcome the deferral option: exterior work is postponable and weather-framed, so the estimate must manufacture urgency (season windows, substrate decay cost) that interior work gets for free from daily visibility. The process difference is the close mechanism, not the lead. [§ 2.3 influencing demand, § 1.1 OrderWinner, § 13 behavioral]

Q6. Which referral channel produces premium-price clients — and why does marketing spend invert that ranking?

A: Past clients, typically: they arrive pre-sold on quality. Spend inverts it because realtor/designer channels have visible, purchasable touchpoints while past-client activation has no line item — so budget follows buyability, not value. Fund the referral loop as a process with an owner. [§ 13 Goodhart, § 1.1 OrderWinner, § 2.3]

Q7. Why do property managers give us turnover repaints but never occupied-unit work?

A: Turnover work is empty, fast, and price-shopped; occupied work risks tenant complaints — they've slotted you as the cheap/volume vendor, a qualifier-level relationship. Winning occupied work requires demonstrating tenant-handling process (scheduling, communication, cleanup), a different service blueprint. [§ 1.1 OrderQualifier, § 8.2 service blueprinting, § 6.3 SERVQUAL assurance]

Q8. Crew retention collapses every October, before the profitable interior season. What do painters know about our winter plan that we haven't said?

A: That there isn't one. Experienced crew leads model your winter from history: hours collapse, so they leave while the leaving is good. Publish the winter capacity plan (interior bookings, guaranteed hours, maintenance projects) — retention is a demand-smoothing problem communicated as an HR problem. [§ 2.1 seasonality, § 7 aggregate planning, § 13]

Q9. Why do cabinets/trim — our highest-margin work — come almost entirely from one estimator's book?

A: Because that estimator sells the craft outcome while others quote wall-square-foot habit. The capability is person-dependent, which means it's uncodified knowledge. Extract their qualification criteria, photo pitch, and pricing frame into standard work before they leave with the margin. [§ 13 standard work/communities of practice, § 1.1 focus]

Q10. Two-coat vs one-coat calls vary between estimators on identical substrates. Whose call is right?

A: The spec's — and you don't have one. Coating system selection should be a function of substrate condition and paint line data sheets, not estimator judgment. Write the decision rule (hide, color change, substrate porosity), then audit both against it; warranty claims will arbitrate. [§ 0.3 specifiedBy, § 6.3 acceptance criteria, D5]

8. Pest control

Q1. Technician turnover spikes exactly when route density should make days easiest. What is route data not showing?

A: The work-content per stop: dense routes often mean smaller accounts, more stops, more doorstep friction — density optimized your drive time, not their day. Measure stops-per-day and per-stop revenue against tenure; if dense = grind, rebalance routes by revenue per stop-hour, not stops per mile. [§ 10 routing, § 4.4 job design, § 13]

Q2. Why do mosquito add-ons attach in some territories and not adjacent ones with identical demographics?

A: Demographics don't buy — micro-geography (standing water, lot size, tree cover) and tech behavior do. Check attach rate by technician within the weak territory first; if it's uniform, it's the land. This is demand heterogeneity that territory averages hide. [§ 2.1 Demand Pattern, § 6.3 stratification, F1]

Q3. Termite inspection-to-treatment conversion varies twofold between identically trained techs. What do high converters say at the kitchen table?

A: They show evidence (mud tubes, damaged wood, thermal images) and quote a monthly number, not a lump sum — converting an abstract risk into a visible, financed decision. Ride along, record, codify; conversion variance that large is standard work waiting to be written. [§ 13 standard work, § 1.1 OrderWinner, § 6.1 perceived quality]

Q4. Which acquisition channel produces customers who stay past year two — and does CAC math know?

A: Almost never: CAC is computed at acquisition, retention arrives years later, and the two are never joined per cohort. Compute channel-level 24-month LTV/CAC; discount channels typically produce deal-seekers with structurally shorter lifetimes. [§ 1.3 metrics, § 13 Goodhart, § 2.1]

Q5. Why do quarterly-plan customers who skip one service cancel at 2.5× — the pest comeback or the broken habit?

A: The broken habit: the service is invisible when it works, so the subscription's value lives in the routine, and one skip breaks the routine's spell. Evidence: cancellations cluster before any pest evidence appears. Engineer the re-engagement within days of a skip. [§ 13 behavioral ops, § 8.2 service blueprint, § 2.3]

Q6. Does per-door pricing assume a treatment time our route data contradicts?

A: Pull actual stop durations from the route app and compare to the standard time inside the price. If actuals exceed the assumption, every dense-route gain is being eaten by underpriced work content — your price book and your process physics are divorced. Reconcile standard times annually. [§ 4.4 work measurement/standard time, § 1.3]

Q7. Why do one-time treatments convert to plans at 15% when the script assumes 40%?

A: Because the script was written from hope, not funnel data, and one-time buyers self-select as problem-solvers, not relationship-buyers. Fix both: re-forecast from actual conversion, and create a distinct follow-up sequence for one-timers instead of the plan pitch at the door. [§ 2.2 forecast bias, § 13, § 2.3]

Q8. Why does our commercial book grow through untraceable referrals?

A: Because B2B referrals travel through property-manager networks with no trackable artifact. Untraceable doesn't mean unmanageable: ask at signing, code it, and the pattern (usually 2–3 connector accounts) emerges. Relationship capital is a resource — put it under measurement. [§ 4.1 Resource taxonomy, § 0.3 measuredBy, § 11.2]

Q9. Are bed bug jobs a profit center or reputation insurance — and do we price accordingly?

A: Decide explicitly. As insurance (defensive work to protect plan accounts), price for cost recovery; as a profit center, price for the expertise premium and warranty risk. The current muddle — premium effort at commodity price — is an ungoverned [Policy](#) gap. [§ 0.3 governedBy, § 1.1 CompetitivePriority, § 6.2]

Q10. When we lose a commercial account, did our own reports flag declining pest activity — did we document ourselves out of the contract?

A: Sometimes, yes: clean reports read as "no pest pressure = no need for service" unless the report narrates prevented activity (bait consumed, entries sealed, seasonal pressure suppressed). The deliverable is risk management, so the report must make the invisible work visible. [§ 6.1 perceived quality, § 8.2 service blueprint, § 2.3]

9. Janitorial / commercial cleaning

Q1. Why do we win school and medical bids but lose office bids at the same price per square foot?

A: Different order winners: schools/medical buy compliance and documented process (your strength); offices buy responsiveness and daytime polish, often decided by the lowest credible price from a relationship vendor. Same service, two products. Segment the bid playbook. [§ 1.1 OrderWinner by segment, § 6.3 SERVQUAL]

Q2. Why do we discover scope creep (extra restrooms, new floors) only at renewal instead of billing monthly?

A: Because no one owns scope verification as an activity — the contract spec is static while the building changes. Institute quarterly scope audits as a scheduled [Activity](#) with an assigned [Resource](#); creep billed monthly is revenue, creep found at renewal is a negotiation. [D5, § 0.3 specifiedBy/governedBy, § 9 change control]

Q3. Lowest-price contracts churn in 14 months; near-lost price deals stay for years. What does the price fight select for?

A: It selects buyers: winning on lowest price acquires procurement-driven clients who'll leave for the next dime; fighting hard and nearly losing means you competed against a relationship they preferred — those clients bought value. Price is a selector of customer type, not just a number. [§ 1.1 OrderWinner, § 2.3, § 13]

Q4. When we absorb supply increases instead of invoking escalation clauses, what are we buying with that margin?

A: Nothing — unilateral goodwill has no accounting identity and no recipient who notices. Escalation clauses exist as [Policy](#); not invoking them is an unpriced transfer. If you want goodwill, invoke the clause and discount visibly; invisible absorption buys zero retention. [§ 0.3 governedBy, § 11.3 price fluctuation, § 13]

Q5. When the client's facility manager changes, why does renewal probability drop regardless of documented performance?

A: Because the relationship was with a person, not an institution — your results were the FM's asset, not the client's. Multi-thread the account (property owner, regional manager) and make performance reporting land with more than one inbox. [§ 11.2 relationship structure, § 13, A8]

Q6. Do day porters retain accounts by adding value or by creating switching costs — and should we care?

A: Care: value-based retention survives a procurement audit; switching-cost retention dies the day a competitor offers transition help. Instrument which it is (satisfaction by porter presence, complaint latency), then price and staff accordingly. [§ 6.3 SERVQUAL responsiveness, § 1.1 dependability, § 2.3]

Q7. Why does per-square-foot margin vary 40% across identical-spec buildings — crew, site manager, or scope drift?

A: Usually scope drift plus crew assignment; identical specs diverge in practice within months. Run a variance decomposition: labor hours per visit × pay rate vs. contract; the building-level labor ledger identifies which factor moves. [§ 6.3 SPC variance decomposition, § 4.4 work measurement, D5]

Q8. Night-crew turnover concentrates in our three most profitable buildings. What can't the P&Ls see?

A: The supervisor behavior at those sites: profitable buildings get the pressure (tightest labor budgets, harshest QC), and night supervision is unobserved. Turnover is a process output; audit the site-level management practice, not the account economics. [§ 13 behavioral ops, § 4.4 job design, § 6.3]

Q9. Which predicts account profitability better: building size or client decision-making structure — and do we score leads on it?

A: Decision-making structure: a single empowered FM yields stable scope and fast issue resolution; committee-managed buildings generate churn and creep. You almost certainly score leads on size alone. Add a decision-structure field to qualification. [§ 1.1 MarketPositioning, § 2.2 causal indicators, § 13]

Q10. Why does inspection score correlate with retention only above a certain contract size?

A: Below that size, the buyer is price-led and inspections are theater; above it, a professional FM actually reads scores and manages by them. Same KPI, two audiences — the KPI only drives outcomes where a [governedBy](#) feedback loop exists. [§ 0.3 measuredBy/governedBy, § 13 Goodhart, § 1.1 segment]

10. Handyman services

Q1. Platform bookings carry 5× our no-show rate at identical job sizes. What does the platform customer owe us that the direct customer does not?

A: Nothing — and that's the answer. The direct customer made a relational commitment (conversation, referral chain); the platform customer made a click. Commitment is an input you can engineer: deposits, confirmation calls, and narrower windows convert platform demand into committed demand. [§ 13 behavioral ops, § 2.3 accommodating demand, § 8.1 abandonment psychology]

Q2. Our best customers eventually hire a specialist directly. Which service is the gateway they stop calling us for first?

A: The highest-stakes one — usually electrical or plumbing work where perceived risk exceeds generalist trust. That's rational: your order winner is convenience for low-risk jobs; specialists own high-risk ones. Decide whether to refer out gracefully (keeping the relationship) or lose the whole customer by fumbling the handoff. [§ 1.1 OrderWinner, § 6.1 perceived quality, § 11.2 make-vs-buy logic]

Q3. Do we lose money on sub-\$200 jobs once drive time is honest — and why do we keep taking them?

A: Allocate drive time and the answer is usually yes. You keep taking them for relationship seeding and schedule-filling — legitimate if priced as marketing, ruinous if priced as revenue. Create a minimum charge or bundle small jobs into route days. [§ 4.1 financial resource, § 10 routing, § 2.3 backlog accommodation]

Q4. Why does materials markup vary by which hardware store is nearest rather than by policy?

A: Because there is no policy — markup is decided per-trip by whoever drives. That's an ungoverned pricing activity producing random margin and occasional customer distrust. Write the price book: markup schedule plus a sourcing rule, and the variance disappears. [§ 0.3 governedBy, § 5.1 consumables, § 13 standard work]

Q5. Which do we chronically underbid: trades-adjacent jobs or odd jobs nobody quotes?

A: Trades-adjacent: you estimate them with handyman heuristics against specialist risk (code, inspection, callbacks) you don't price. Odd jobs are priced by pure judgment but carry no hidden tail. Audit margins by job class; the regulated-adjacent work will show the leak. [§ 9 risk, § 6.2 external failure, § 4.4 work measurement]

Q6. Average ticket grows yearly while jobs-per-day falls — emerging strategy or failing scheduling?

A: Check the cause: if bigger jobs are chosen (mix shift), it's emergent strategy — codify it. If drive time and gaps are growing, scheduling is failing and ticket size is just mix noise. Decompose jobs-per-day into drive, gaps, and on-site hours. [§ 1.1 deliberate vs emergent strategy, § 7 scheduling, F1]

Q7. What percentage of "I'll text you a price tonight" jobs do we win — and why do we keep doing it?

A: Track it: typically under half the on-the-spot quote rate, because evening quotes land after the customer's urgency cooled and competitors quoted live. You keep doing it to protect the day's schedule. If the win-rate gap is real, the schedule is eating the pipeline. [§ 1.1 OrderWinner speed, § 2.3, § 8.1]

Q8. Customers ask for half-day/full-day packages; we persist hourly. Packaging failure or pricing fear?

A: Pricing fear — hourly feels safe because it transfers overrun risk to the customer, but customers read hourly as open-ended cost anxiety. Day-rate packages convert overrun risk into a premium you can price. Test one packaged SKU against hourly on matched jobs. [§ 2.3 pricing/demand shaping, § 6.1 perceived quality, § 13]

Q9. Adding a second tech halved margin instead of doubling revenue. Which constraint did we hit: sales, scheduling, or supervision?

A: Diagnose by the invariant: if the pipeline didn't double, you hit demand (A2 — throughput is capped by $\min(\text{capacity}, \text{demand})$, and you elevated the wrong constraint). If jobs were there but days fragmented, scheduling. Post-hoc, the P&L split will show idle hours vs. missing jobs. [A2, A3, § 16 TOC exploit → subordinate → elevate]

Q10. When a customer becomes a repeat, what was the first job — and should we engineer more of those?

A: Pull the cohort data: repeat customers almost always trace to a first job that was urgent, small, and flawlessly executed — the trust trial. If so, treat that job class as your acquisition product: price it to win, staff it with your best, and instrument the 90-day callback rate. [§ 2.2 cohort analysis, § 1.1 OrderWinner, § 6.3 SERVQUAL reliability]

11. Flooring contractors

Q1. Hardwood refinishing has our highest margins and lowest review volume; carpet does the reverse. Which does the marketing calendar feature?

A: Carpet, almost certainly — marketing follows lead volume, not contribution. Reviews skew to carpet because it's a bigger customer pool; margin lives in refinishing. Rebalance the calendar by margin-per-lead, and build a review-harvest step into refinishing closeout. [§ 1.3 metrics, § 13 Goodhart, § 2.3]

Q2. Insurance water-damage jobs arrive through two specific plumbers. What happens if either retires?

A: A single-point-of-failure in your demand network: two nodes carry a channel. Series reliability logic — quantify the revenue at risk, then deliberately widen the referral base (more plumbers, adjusters, mitigation firms) before the failure event, not after. [§ 4.3 series reliability, § 11.4 risk assessment/mitigation]

Q3. Our LVP price sits 15% above market and close rate hasn't moved. Which assumption does that break?

A: The assumption that you're in a price-shopped market. Flat close rate above market means your customers buy trust/installation quality — LVP is a qualifier, not the winner. But test the ceiling: the flat curve means you've never found your price boundary. [§ 1.1 OrderQualifier vs OrderWinner, § 2.3 pricing, § 6.1]

Q4. Does steady spec-home builder work build capacity or slowly replace a profitable retail book with a captive one?

A: Watch the mix trend: builder work is dependable but price-capped and payment-slow; if its share creeps past ~50%, your crews' skills, schedule, and cash cycle reconfigure around one buyer — a supply-chain dependence you chose by default. Cap the share deliberately. [§ 11.2 sourcing structure/single-source risk, § 1.1 VerticalIntegration decisions, § 4.1 working capital]

Q5. Why do moisture failures cluster where we subcontracted the prep rather than the install?

A: Because moisture mitigation is the spec-critical hidden layer, and subs optimize for speed on invisible work you can't inspect at handover. The failure is yours at warranty time. Either in-house the prep or make moisture readings a documented acceptance gate with the sub. [§ 6.3 acceptance sampling/criteria, § 11.2 supplier management, A8]

Q6. Does the showroom convert browsers or validate decisions made online — and should we know?

A: Yes, you should: the two roles have opposite implications. If it converts, invest in display and staff; if it validates, the showroom is a cost center supporting an online-led funnel and should shrink or become appointment-only. Instrument the entry question ("have you chosen a product?"). [§ 10 facility/fixed-position layout economics, § 2.3, § 8.2 blueprint]

Q7. Why does material waste double on jobs where the client supplies the flooring?

A: Two causes: client-bought material comes in wrong quantities/quality (more cutting around defects), and your crews treat client stock as free — no skin in the waste. Add a waste-allowance clause for client-supplied material and inspect on receipt. [§ 5.1 Material role, § 6.2 internal failure, § 0.3 governedBy]

Q8. Would dropping the lowest-margin product line raise total profit — and why does nobody want to run the test?

A: Probably yes if freed capacity migrates to higher-margin work; no if the line feeds loss-leader traffic. Nobody runs it because the line's overhead absorption is invisible and killing it feels like shrinking. Model contribution margin after reallocating crew-hours, then decide. [§ 1.3 productivity, § 1.1 focus, A2]

Q9. When clients delay after material delivery, who absorbs storage and double-handling — and why isn't it in the contract?

A: You do, because the contract treats delivery as a milestone without a delay contingency. Conservation of flow: the inventory has to sit somewhere, and it's sitting on your working capital. Add a delay/storage clause; double-handling is pure [NVA](#) you're currently donating. [A5 conservation of flow, D3 NVA, § 0.3 governedBy]

Q10. Estimators' waste factors range 7–15% on identical layouts. Whose number is right?

A: Historical job data is right: pull actual material ordered vs. installed per layout type and set the standard waste factor by product and room geometry. A 7–15% spread means estimating is opinion, not measurement — the spread itself is costing you either margin (high) or jobs (low). [§ 4.4 work measurement, § 13 standard work, § 5.1]

12. Tree services

Q1. Customers quoted in winter hire whoever answers first in spring. What would a March follow-up sequence be worth?

A: Compute it: winter quote volume × historical close rate × recovered share × margin. Dormant quotes are WIP in your sales process; a two-touch March sequence typically recovers a third of them at near-zero cost — the cheapest capacity you own. [§ 5.1 WIP analogy, § 2.3 accommodating demand, § 1.3]

Q2. Which predicts a profitable month better: booked revenue or job mix within a 15-minute drive?

A: Drive-time mix. Your binding constraint is crew-day capacity, and drive time consumes it without revenue; booked revenue ignores geography. Test: regress monthly margin on both — clustered revenue beats scattered revenue every time in field services. [A2/A3 constraint logic, § 10 routing/location models, § 1.3]

Q3. Does crane work subsidize climbing crews or vice versa — can job costing separate them?

A: Only if equipment ownership cost is allocated per crane-day, which most job costing never does (crane depreciation sits in overhead). Allocate true crane cost per deployed day and re-run margins; usually the crane wins on technical jobs and loses on jobs where it was used for convenience. [§ 4.1 Physical resource costing, § 1.3, D5]

Q4. Why does stump-grinding attach rate vary by crew lead rather than customer segment?

A: Because attachment is a sales behavior, not a demand property: some leads present the stump as an eyesore-resolution with a same-day discount, others never mention it. Segment-independence proves the variance is process, not market. Codify the offer. [§ 2.3 demand shaping, § 13 standard work, § 6.3 stratification]

Q5. Why quote per tree for removals but per hour for pruning — and does either reflect how work goes?

A: Per-tree works because removal time scales with tree size class (roughly standardized); per-hour for pruning confesses you've never built a pruning work standard. Build unit standards (per tree size/species) and convert pruning to fixed quotes with defined scope — customers buy certainty. [§ 4.4 work measurement, § 2.3 pricing, § 6.1]

Q6. Municipal contracts auto-renew while commercial re-bids annually at lower prices. What does the municipal relationship have?

A: Switching costs plus a risk-averse buyer: municipalities pay for dependability and paperwork compliance; commercial PMs are bonused on visible savings. You're a [strategic](#) supplier in one Kraljic quadrant and a [leverage](#) item in the other. Bid commercial on risk-transfer, not price. [§ 11.2 Kraljic matrix, § 1.1 OrderWinner by segment]

Q7. Reviews praise "showed up when they said," never price — while marketing leads with affordability. Which promise should the ads make?

A: Reliability. The review corpus is the voice-of-customer telling you the actual order winner: in a market of no-shows, kept appointments win. Reposition ads around scheduling certainty ("crews arrive on the booked day") and let price be the qualifier. [§ 6.3 QFD voice-of-customer, § 1.1 OrderWinner, § 6.3 SERVQUAL reliability]

Q8. When we decline a hazardous removal, does the customer call a less insured competitor — and what is our liability strategy optimizing?

A: Yes, they do. Your strategy optimizes tail-risk avoidance, which is rational, but unexamined: define the hazard classes you decline, price the ones you accept to carry the risk premium, and consider a referral agreement with a specialist. Risk management is a priced product, not a reflex. [§ 11.4 risk assessment, § 1.1 CompetitivePriority, § 2.3]

Q9. Why do emergency storm calls produce our best revenue days and worst safety incidents — connected by scheduling?

A: Yes: storm surges push utilization past 1 with fatigued crews on degraded sites — the variability–buffer law applied to safety. Incidents concentrate when schedule pressure meets overtime. Storm pricing should buy you the buffer (more crews, mandated rest), not just more revenue. [A4, § 14 safety management, § 4.4 ergonomics/fatigue]

Q10. When a competitor undercuts 30% on a removal — uninsured, or do they know something about disposal costs?

A: Check disposal first: dump fees and chip/log resale are the hidden margin lever in tree work, and a competitor with a disposal arrangement can legitimately undercut 20–30%. If their certificates check out, audit your own disposal line before assuming recklessness. [§ 11.2 supply network, § 4.1 materials, § 1.3 cost analysis]

13. Pressure washing

Q1. We lose fleet-washing bids to companies with worse equipment and dedicated account managers. What are fleet buyers purchasing?

A: Administrative relief, not clean trucks: consolidated invoicing, one throat to choke, scheduled compliance. The account manager is the product; equipment is a qualifier. Productize the account service layer (reporting, scheduling, single invoice) and re-bid. [§ 1.1 OrderWinner, § 6.3 SERVQUAL assurance, § 8.2 blueprint]

Q2. Which business do our systems assume we run: residential with commercial jobs, or commercial with residential distractions?

A: Audit the systems: invoicing cadence, scheduling horizon, and marketing spend will confess one identity while revenue mixes two. Mixed operations without explicit segmentation underperform a focused one — pick the primary and make the other a deliberate, ring-fenced secondary. [§ 1.2

Focus/plant-within-a-plant, § 1.1 StructuralDecisions]

Q3. House-wash customers rebook at 14 months, driveway at 22. Does the reminder system know?

A: It should and almost certainly doesn't — most reminder systems run one cadence. Segment rebooking cycles by service type (per [FlowUnit](#) class); a reminder fired at month 12 for driveways is noise, at month 18 it's revenue. [§ 2.1 Demand Pattern by segment, § 12 platform layer, § 2.3]

Q4. Does the customer perceive a soft-wash vs pressure-wash difference worth the price gap?

A: Only when framed as risk: "won't strip paint/damage siding" converts the technical distinction into perceived quality. If you sell it as a method, it's invisible; sell it as asset protection and the gap justifies itself. Test both scripts on matched quotes. [§ 6.1 perceived quality/durability, § 2.3, § 13]

Q5. Why do HOA contracts generate near-zero referrals to homeowners in the same community?

A: Because the HOA buyer (board) and the homeowner are different flow units with no designed handoff: the crew works common areas invisibly, and no artifact reaches the homeowner. Engineer the transfer — door hangers after common-area work, community pricing days. [§ 8.2 service blueprint/line of visibility, § 2.3, § 13]

Q6. Fuel and chemical costs rise 20%: we reprice commercial immediately but hold residential flat for a year. Why?

A: Contractual cover: commercial agreements have escalation clauses, residential is sold as a one-time price you're afraid to move. But residential rebooking is habitual, not contractual — a modest increase with notice rarely breaks the habit. You're donating margin to an untested fear.

[§ 11.3 price fluctuation, § 0.3 governedBy, § 2.3]

Q7. Do seasonal prepaid packages smooth cash flow or discount our best months? What does redemption timing show?

A: Pull redemption curves: if redemptions cluster in your peak months, you sold your scarcest capacity at a discount — anti-yield-management. Restructure so prepaid redemptions are limited to shoulder months, converting the package into a demand-smoothing tool. [§ 8.2 yield

management/fences, § 2.3 accommodating demand, § 4.1 working capital]

Q8. Instant online quotes close below phone quotes at identical pricing. What does the phone add?

A: Trust construction and objection handling — the assurance dimension. The form gives a number; the conversation gives confidence the number is right for their surfaces. Add photo/video upload plus a callback option to approximate the conversation. [§ 6.3 SERVQUAL assurance, § 8.2 service

blueprint, § 12]

Q9. Why does per-job margin fall when we book two jobs on the same street — the opposite of route density?

A: Something else moves with same-street bookings: usually discounting ("neighbor discount") or longer combined setup (one water source, parking constraints) eating the saved drive time. Decompose the margin fall into price vs. time — density math only works if price and work content hold. [§ 10 routing, § 1.3, F1]

Q10. Why does revenue per tech-day vary 50% between two crews with same equipment and routes?

A: With inputs held constant, the variance is human: upsell behavior, pace, and scope interpretation at the doorstep. Ride both crews, measure per-stop revenue and duration, and codify the high crew's doorstep routine into standard work. [§ 13 standard work, § 4.4 work measurement, § 6.3 special cause]

14. Water/fire/mold restoration

Q1. Equipment utilization has never exceeded 60%, yet we keep buying air movers. What approved the last purchase?

A: Job-level panic, not capacity analysis: each loss site demands equipment per IICRC class, and purchases are justified per-incident instead of against fleet utilization. The decision rule should be utilization-at-peak, not utilization-on-average — but 60% average with weekly rentals means the fleet mix is wrong, not small. [§ 4.2 design vs effective capacity, § 1.1 StructuralDecisions/Capacity, § 5.1]

Q2. When we lose a contents pack-out to the homeowner's DIY decision, what did the estimator say in the first ten minutes?

A: Probably led with process and price instead of risk: DIY pack-outs fail on smoke residue, mold cross-contamination, and insurance documentation. The first ten minutes must convert the pack-out from a cost into claim protection — that's the order winner. Script it. [§ 1.1 OrderWinner, § 6.3 SERVQUAL assurance, § 13]

Q3. Does deploying to a hurricane three states away strengthen or hollow out home-market response capacity?

A: Both, in sequence: CAT deployment is high-margin but converts your home capacity to zero — if a local loss event hits during deployment, your referral sources (plumbers, agents) experience a stockout and reroute permanently. Set a home-market reserve capacity rule before chasing storms. [§ 11.4 resilience/TTR vs TTS, A4 buffer choice, § 1.1 StructuralDecisions]

Q4. PMs' cycle times improve right before their quality scores fall. What is the bonus plan rewarding?

A: Speed at the expense of drying verification — Goodhart in pure form: cycle time is the measured proxy, quality (moisture mapping, documentation) is the casualty. Rebalance the scorecard to include re-dry rates and supplement quality, or the metric will keep eating the outcome. [§ 13 Goodhart's law, § 1.3 balanced metrics, § 6.3 SPC]

Q5. Why do we know insurance close rate to the decimal but never measured direct-pay retail win rate?

A: Because insurance volume arrives through structured channels with built-in tracking, while retail losses are ad hoc — so measurement followed the plumbing, not the strategy. Retail losses are your margin; instrument that funnel (source, quote, win) and it will likely out-convert the program work you over-manage. [§ 0.3 measuredBy, § 13 Goodhart, § 1.1 MarketPositioning]

Q6. Why does average job size drop when we add a service territory?

A: Dilution: new territories add referral sources at the margin (smaller plumbers, distant agents) whose losses skew smaller, and your response time to large losses degrades with distance. Segment job size by territory age; the fix is channel development in the new territory, not more territory. [§ 10 location/coverage models, § 2.1 demand mix, § 11.2]

Q7. Which costs more: TPA program discounts or the marketing to replace that volume with direct jobs?

A: Compute replacement cost honestly: TPA discount % $\times$ program revenue vs. CAC $\times$ direct jobs needed. TPAs typically cost 10–15% off the top plus compliance overhead; direct replacement via plumber/agent channels usually wins at scale — but only if you fund the channel like a sales force, not a brochure. [§ 11.2 supplier/channel economics, § 1.3, § 2.3]

Q8. Mitigation invoices get negotiated down by adjuster more than by documentation quality. What are we pricing: the work or the adjuster?

A: The adjuster — your realized rate is adjuster-specific, which means the negotiation process, not the work standard, sets price. Counter with documentation-as-spec: Xactimate-line photo evidence per line item moves the negotiation from relationship to record. [§ 0.3 specifiedBy, § 11.3 coordination, § 13 standard work]

Q9. Why do plumber referrals produce smaller losses than agent referrals — and should relationship spend reflect it?

A: Yes, but measure job value, not count: plumbers see supply-line leaks early (small water jobs); agents see fires and major claims (large, rare). Allocate relationship budget by expected margin per referral, and you'll likely shift spend toward agents while keeping plumbers for volume. [§ 2.2 channel analysis, § 1.3, § 11.2]

Q10. Do reconstruction upsells profit from the mitigation relationship, or does bundling make us slower to approve than specialists?

A: Test approval latency: if your reconstruction bids take longer because mitigation documentation competes for the PM's attention, the bundle costs you the rebuild. The ontology answer: the mitigation-to-rebuild handoff is a decoupling point — manage it as one, with a dedicated rebuild estimator. [§ 1.2 DecouplingPoint, § 9 project management, § 1.1 focus]

15. Pool service & repair

Q1. Repair margins shrink as the route book grows, though scale should cut parts costs.

Where does scale cost us?

A: In complexity: more accounts mean more equipment variety (SKUs of valves, boards, heaters), so techs arrive without the right part more often — second trips eat the parts discount. Counter with van-stock standardization by route mix and a first-visit-fix KPI. [§ 5.1 MRO/safety stock, § 5.2 ABC, A4]

Q2. Green-to-clean jobs convert to weekly service at 10% despite obvious recurring need.

Where does the rescue-to-relationship handoff break?

A: At the moment of triumph: the green-to-clean delivers a visible miracle, the customer feels "done," and nobody converts the relief into a maintenance contract on the spot. The handoff is an activity with no owner and no script — assign it, price the transition (first month free with 12-month), measure it. [D5, § 2.3, § 8.2 service blueprint]

Q3. Are upgrade recommendations driven by tech judgment or supplier rebate tiers — and could a customer tell?

A: Check the data: brand mix of recommended equipment vs. rebate schedule will answer instantly. If rebates drive it, you're running an undisclosed conflict that one informed customer can expose — and in a referral business, that's an external-failure reputational risk. Disclose or decouple. [§ 11.3 incentive misalignment, § 6.2 external failure, § 13]

Q4. Why do winter-paused customers restart with competitors at twice the rate of winter-pricing customers?

A: Out of sight, out of contract: pausing severs the relationship artifact (the monthly visit/invoice) and spring restart becomes an open auction. Winter pricing keeps the thread unbroken — it's cheap retention insurance. The pause option is a churn machine wearing a customer-service costume. [§ 2.3 demand shaping, § 13 behavioral, § 8.2]

Q5. When a freeze damages 200 pools at once, which accounts do we triage first — and is that policy written?

A: It isn't, and it should be: triage by customer lifetime value, vulnerability (screen enclosures, automation at risk), and referral influence — a pre-committed [Policy](#) beats ad-hoc heroics, and customers forgive delays they were told to expect. Write the surge protocol before the surge. [§ 11.4 BCP, § 0.3 governedBy, § 8.2 yield/allocation]

Q6. Why do we win new-build pools for two years then lose them en masse to a low-price entrant?

A: Because nothing accumulates: after the builder handoff and warranty period, you're an undifferentiated monthly charge, and the neighborhood buys as a social block — one defection cascades. Build lock-in before year two: equipment data, loyalty pricing, neighborhood presence.

[§ 1.1 OrderQualifier migration, § 11.2 switching costs, § 13 network effects]

Q7. Does chemical-inclusive pricing protect margin or hide cost creep? What did last summer do?

A: Audit a sample: actual chemical cost per pool vs. the flat allocation. Chemical-inclusive pricing is a buffer that works until usage drifts (heat waves, algae blooms) — without a variance check, creep hides for years. Reconcile quarterly and add an extreme-condition surcharge clause. [§ 5.1 consumables, § 1.3 variance, § 0.3 governedBy]

Q8. When a tech quits, why do route accounts churn at double baseline even when service quality doesn't measurably change?

A: Because the account asset was the tech's relationship and presence, not the service record — customers churn on continuity loss, not quality loss. Institutionalize the relationship: route notes, customer preferences, and a warm handoff visit from the replacement. [§ 13 knowledge/relationship capital, § 6.3 SERVQUAL empathy, § 8.2]

Q9. Which accounts lose money once drive time is allocated honestly — and why is route optimization always postponed?

A: The rural outliers — always. Postponement happens because dropping accounts feels like shrinking and optimization tools cost money, while the leak is diffuse. Run drive-time-allocated margin per account once; the bottom 5–10% usually funds the entire optimization project. [§ 10 routing/location models, § 1.3, § 4.1 working capital]

Q10. Why does our renovation/resurface work come from other companies' service accounts more than our own?

A: Because your own customers see you as the weekly-chemical guy — your category anchor blocks the bigger purchase. Competitors' unhappy accounts shop around for renovation and find you. Fix: surface renovation capability to your base (equipment-age-triggered proposals) before the anchor sets. [§ 1.1 MarketPositioning, § 6.1 perceived quality, § 2.3]

16. Concrete / masonry

Q1. Flatwork wins at 40%, decorative at 70% under the same pricing logic. What does the decorative buyer hear?

A: A design sale — decorative quotes come with samples, pictures, and transformation language; flatwork quotes are commodities priced per square foot against three competitors. The flatwork buyer hears a price; the decorative buyer hears a vision. Sell flatwork on base prep and warranty (the invisible quality) to escape the commodity frame. [§ 1.1 OrderWinner, § 6.1 perceived quality/aesthetics, § 2.3]

Q2. When a homeowner chooses pavers over poured concrete, did we lose on price or fail the maintenance case?

A: Usually the maintenance case was never made: pavers win the "repairable, no cracks" narrative by default. Concrete's counter (cost per year, no weed joints, modern finishes) requires a comparison sales asset you apparently don't present. Build the side-by-side. [§ 1.1 OrderWinner, § 6.3 QFD, § 2.3]

Q3. When a GC squeezes 8% at award, where does the margin come out: mix design, labor pacing, or profit?

A: Profit — if it came out of mix design or pacing, you'd be buying callbacks (external failure) to fund a discount. That's the test: if quality inputs are untouched, the squeeze was affordable; if crews rush or mixes lean out, the 8% returns as warranty with interest. [§ 6.2 CoQ trade-offs, § 11.3 coordination failure, § 13]

Q4. Why does our winter revenue plan assume indoor work we've never sold?

A: Because the plan is a wish, not a demand model: indoor concrete (garage coatings, basements) is a real counter-seasonal product, but it needs its own channel and sales motion. Either build that product deliberately this off-season or remove it from the plan and manage winter as a capacity problem. [§ 2.3 counter-seasonal strategy, § 7 aggregate planning, A10]

Q5. Do decorative finishes attract craft-valuing clients or change-order artists? What does the change-order log say?

A: Read the log: if mid-pour changes cluster on decorative jobs, the segment buys a vision they refine live — which is fine only if change orders are priced and scheduled as rework events. Unpriced mid-pour changes are the most expensive [NVA](#) in construction. [D3, § 9 change control, § 2.1 demand type]

Q6. Why does pump-truck scheduling set our job calendar more than the sales pipeline?

A: Because the pump is your true bottleneck resource — a shared, scarce, expensive asset. The ontology is blunt: throughput $\leq$ bottleneck capacity, so the bottleneck's calendar legitimately governs. Own more pump capacity, contract priority windows, or sequence pours to minimize pump moves. [A3, A2, § 7 finite-capacity scheduling]

Q7. Why do cracking callbacks cluster in our busiest month rather than our coldest?

A: Busy-month pours get compressed curing attention: crews strip forms early, skip curing compound, and stack loads on green slabs to keep schedule. Cold weather you manage deliberately; schedule pressure you don't. Track callbacks by pour-week crew load, not by temperature. [§ 6.3 SPC stratification, A4, § 6.2 external failure]

Q8. Why does material cost % of revenue drift up yearly despite locked supplier pricing?

A: Because the drift isn't price — it's waste, over-pour, and mix escalation (customers spec'ing higher PSI), plus short-load fees on bad estimates. Supplier price is one term; yield is the other. Measure yield variance (ordered vs. theoretically required) per job. [§ 5.1 materials, § 1.3 variance analysis, § 4.4]

Q9. We keep bidding 90-day-pay subdivision work while complaining about cash. What if cash conversion set the bid criteria?

A: Then half that work would fail the screen. Working capital is a [Resource](#); a job that ties cash for 90 days must earn a financing premium. Add a cash-conversion score (days to pay $\times$ margin) to bid/no-bid, and watch the calendar recompose toward retail and smaller GCs. [§ 4.1 Financial resource, § 1.1 StructuralDecisions, § 11.2]

Q10. Which crew lead's jobs come back for warranty least — and have we studied what they do at the pour?

A: That's the study to run: warranty rate by crew lead identifies your internal best practice. In concrete it usually traces to base prep patience and curing discipline. Video it, codify it, make it the standard work — the gap between your best and average lead is your cheapest quality program.

[§ 13 standard work/communities of practice, § 6.2 prevention, § 6.3]

17. Fencing contractors

Q1. Why do leaning-post warranty calls trace to soil type more than installation method — and does the quote account for soil?

A: Because soil governs post stability and your quote treats setting as uniform. Add a soil/condition variable to estimating (depth, concrete volume, post spacing adjustments) — the warranty data gives you the map. An unpriced input variable is a margin leak with a lag. [§ 0.3 specifiedBy, § 6.2 external failure, § 2.2 causal]

Q2. Which is the real business: installation, or permit-and-survey navigation — and do we price like we know?

A: For most residential buyers, the navigation is the differentiator — any crew can set posts; only you can make the paperwork, line disputes, and utility locates vanish. If permits consume hours but appear nowhere on the quote, you're giving away the actual product. Price it as a line item.

[§ 1.1 OrderWinner, § 8.2 service blueprint, D3]

Q3. Our post-hole subcontractor's availability sets our schedule. At what volume does in-housing pay?

A: When $(\text{sub premium} + \text{schedule-loss cost}) \times \text{annual volume}$ exceeds the loaded cost of equipment + operator at realistic utilization. Include the shadow price of the schedule control you lose — a bottleneck you don't own still caps your throughput. Run the make-vs-buy with delay cost in it. [§ 11.2 make-vs-buy, A3, § 4.2]

Q4. Shared-property-line jobs finish fastest and generate the most disputes. What does speed cost on a shared line?

A: Consent verification. Fast completion skips the slow step — documented neighbor agreement on line location and cost share — and the dispute arrives post-hoc when positions harden. Make neighbor sign-off a hard predecessor activity; the day you lose at the start costs less than the week you lose at the end. [§ 3.2 predecessors, § 9 CPM, § 6.2 external failure]

Q5. Why do wood quotes convert immediately while vinyl stalls for weeks, despite better margin?

A: Wood is a known reference price; vinyl requires a value argument (lifetime cost, no maintenance) your quote makes verbally if at all. The stall is the customer doing research your quote should have done. Add a 10-year cost comparison to the vinyl quote and watch the stall shrink. [§ 2.3, § 6.1 durability/perceived quality, § 1.1]

Q6. When a neighbor splits a fence cost, why do we treat it as one sale instead of two future customers?

A: Process blindness: your CRM keys on the payer, so the neighbor — who just watched your crew work for days — vanishes as a lead. Capture both parties as contacts; the neighbor's own fence needs are now warm, and the second sale costs almost nothing. [§ 2.3, § 12 platform/CRM, § 13]

Q7. Do gate/automation add-ons attach because customers want them or because one salesman offers them? What when he leaves?

A: Test by attaching the offer to another rep's script for a month. If attachment follows the person, the capability is uncodified and your margin walks out with him. If it follows the offer, institutionalize it. Either way, the experiment is cheap and the answer is structural. [§ 13 standard work, § 6.3 stratification, § 11.4 key-person risk]

Q8. Wood line moves with lumber market while vinyl tags haven't changed in two years despite supplier increases. What margin are we donating?

A: The full two-year cumulative supplier increase — compute it from invoices. The asymmetry (reacting to lumber, frozen on vinyl) is anchoring: wood feels commodity, vinyl feels fixed-price. Index both lines to supplier cost with a quarterly reprice policy. [§ 11.3 price fluctuation, § 0.3 governedBy, § 13 anchoring]

Q9. Why does material waste on racked terrain exceed the estimating table — and who owns the table?

A: Because the table assumes flat lots; racked panels on slopes cut differently and waste non-linearly with grade. Nobody owns the table — that's the second answer. Assign estimating-table stewardship and add a slope factor; warranty of the estimate is a maintained spec, not folklore. [§ 4.4 work measurement, § 0.3 specifiedBy, D5]

Q10. Why do we lose HOA/commercial bids at prices where we win residential — what's the buyer evaluating?

A: Institutional credibility: insurance certificates, crew size for deadline, warranty backing, references from other HOAs, and bid-document polish. Residential buys price-plus-trust; commercial buys risk transfer. Assemble the commercial bid package as a product, or stay out of that segment deliberately. [§ 1.1 OrderQualifier, § 6.3 SERVQUAL assurance, § 11.2]

18. Garage door services

Q1. Why does builder warranty service cost more than the install margin on new-construction homes?

A: Because builder-grade installs (their spec, their schedule pressure) seed 12 months of adjustments you service for free, and the homeowner relationship belongs to the builder. Price the warranty tail into builder quotes, or exit: revenue that costs more than it earns is negative-margin volume. [§ 6.2 external failure, § 1.3, § 11.3 double marginalization]

Q2. Opener installs: higher satisfaction, lower revenue per hour than repairs. Which does the schedule favor?

A: Probably repairs (urgency wins dispatch), which is correct for margin — but openers build the review base and future replacement pipeline. Run both deliberately: repairs for today's cash, openers for tomorrow's demand asset, and don't let urgency starve the strategic product. [§ 7 dispatch rules, § 1.1 CompetitivePriority, § 2.3]

Q3. When a customer declines the safety-sensor upgrade, did the tech explain liability or just price — and do we record which?

A: You don't record it, so you can't improve it — the decline reason is the most valuable field in the interaction. Techs who explain liability ("your door can kill a pet without these") convert; those who quote a price don't. Log decline reasons, train on the delta. [§ 0.3 measuredBy, § 13 standard work, § 6.1 reliability]

Q4. Spring markup and opener markup differ with no cost rationale. What would a customer find comparing our price book to the parts catalog?

A: An inconsistency that reads as opportunism: springs carry emergency markup, openers carry retail markup, and neither matches cost logic. In an era of phone-price-checks, irrational price books are reputational risk. Rebuild the book on a stated markup policy by category. [§ 0.3 governedBy, § 6.1 perceived quality, § 2.3]

Q5. Spring calls on visibly 20-year-old doors convert to full-door upgrades at 8%. What happens in that conversation — and what should?

A: The tech fixes the spring and leaves — order-taking, not advising. What should happen: a 90-second total-cost comparison (spring now + panels soon + opener eventually vs. full door today) with photos. At 20 years, the upgrade is usually the honest recommendation; 8% means it's rarely made. [§ 2.3, § 6.3 SERVQUAL assurance, § 13]

Q6. Why does same-day service hold in the morning and collapse after 2 PM — staffing curve or dispatch discipline?

A: Test: if morning calls get same-day slots while afternoon calls get tomorrow regardless of truck availability, it's discipline (dispatch stops promising); if trucks are genuinely full by 2 PM, it's a capacity curve mismatched to a demand curve that peaks midday. Stagger shift starts to the demand curve. [§ 8.1 queuing/arrival patterns, § 7 scheduling, § 2.1]

Q7. Why does one tech generate 80% of five-star reviews while two others generate all the callbacks?

A: Because review generation is a behavior (asking + explaining + clean work) and callbacks are a behavior (rushed diagnosis), and you've never split the two KPIs by tech. Both are codifiable: the star tech's close-out routine and the callback techs' failure modes. This is special-cause management, not luck. [§ 6.3 SPC special cause, § 13 standard work, § 0.3 measuredBy]

Q8. Are evening emergency calls our most profitable work or our most expensive marketing, once overtime is honest?

A: Price it and find out: loaded evening cost (OT + callback risk on fatigued work) vs. the emergency premium. Usually profitable per-call but corrosive via next-day capacity debt and error rates. If you keep it, price it explicitly as a premium tier rather than an implicit expectation. [§ 8.2 yield management, § 4.4 fatigue, § 6.2]

Q9. Do we lose commercial overhead-door contracts on price or because we never show a preventive-maintenance plan?

A: The missing PM plan: commercial buyers need to justify the contract to procurement as downtime-risk reduction. A competitor with a documented PM schedule wins the risk-transfer argument even at higher price. Productize the PM plan (inspection cadence, response SLA, documentation) and lead with it. [§ 4.3 preventive maintenance, § 1.1 OrderWinner, § 6.3 assurance]

Q10. Why do property managers churn us every two years — and does our invoicing cycle make us look interchangeable?

A: Yes: if the only artifact they see is a monthly invoice identical in form to every vendor's, you've supplied no differentiation data. PMs churn vendors who give procurement nothing to defend. Send a quarterly service summary (doors serviced, failures prevented, response times) — make the relationship legible. [§ 6.1 perceived quality, § 8.2, § 11.2 supplier scorecards in reverse]

19. Appliance repair

Q1. Why approve \$200 on a \$400 appliance but decline \$300 on a \$1,200 one?

A: Mental accounting, not arithmetic: the cheap appliance reads as "keep it alive," the expensive one reads as "protect my investment from a failing asset" — plus replacement anchoring differs. The repair-vs-replace conversation needs a decision rule (50% of replacement cost, age-adjusted) presented as expertise, not a number. [§ 13 behavioral ops, § 2.3, § 6.3 assurance]

Q2. Why measure tech productivity in calls per day when revenue per van-hour ranks them differently?

A: Because calls-per-day is countable and van-hour economics requires allocation you've avoided. The metric picks the behavior: calls-per-day breeds rushed diagnostics and callbacks; revenue per van-hour rewards first-visit fixes. Switch the primary KPI and watch behavior migrate. [§ 13 Goodhart, § 0.3 measuredBy, § 1.3 productivity]

Q3. Which margin comes from parts markup customers can now check in the driveway — and what's the plan?

A: Quantify the exposure: parts markup share of gross margin. The plan: reframe price as total-cost-of-repair (diagnosis + part + labor + warranty) with the warranty and same-day availability as the justification for markup. Customers accept premiums for speed and guarantee; they punish unexplained ones. [§ 2.3, § 6.1 perceived quality, § 11.2]

Q4. Why do sealed-system jobs sit at the bottom of every tech's preference list despite top billing rates?

A: Because billing rate $\neq$ tech economics: sealed-system work is long, high-variance, callback-prone, and shop-equipment-dependent. The rate premium doesn't compensate the risk premium. Either price it to its true variance or refer it out deliberately — avoided work is a queue distortion. [§ 2.1 process variability, § 6.2 internal failure, § 11.2]

Q5. First-visit fix rates vary 25% between techs with same brands and van stock. What's in the high fixers' vans?

A: Not parts — pre-visit triage: high fixers query symptoms at booking and pre-load the van for the probable failure. The difference is an information activity upstream of the visit. Codify the triage questionnaire into dispatch and the rate converges. [§ 8.2 service blueprint, § 13 standard work, § 5.1 van stock policy]

Q6. Why keep diagnosing appliances we know will be replaced — and is the diagnostic fee compensating honestly?

A: Because the fee funds the trip. If the diagnostic fee covers fully-loaded cost, it's honest triage work with value (the replace decision is worth paying for); if it's a loss-leader hoping for repair authorization, it's a subsidy with bad conversion. State the fee as "decision service" and the ethics resolve. [§ 2.3, D3 VA definition, § 6.3 assurance]

Q7. Do all-day windows cost more in cancellations than 2-hour windows save in routing?

A: Usually yes: all-day windows generate morning-of cancellations and phone anxiety (customers plan around you all day), while 2-hour windows cost routing slack. Measure cancellation rate by window type; the industry answer has converged on narrow windows with live tracking for a reason. [§ 8.1 queue abandonment psychology, § 10 routing, § 2.3]

Q8. A competitor declares the unit dead; the customer calls us for a second opinion. What's our win rate — and does marketing know this moment exists?

A: Measure it — second-opinion calls are your highest-trust inbound and typically convert to repair authorization or replacement referrals at exceptional rates. If marketing doesn't tag this source, your best channel is invisible. Track, name it ("second-opinion service"), and encourage reviews mentioning it. [§ 2.2 channel measurement, § 6.3 assurance, § 13]

Q9. Why do repeat customers generate less revenue than referral data predicts — where do they go between appliances?

A: To whoever is visible at the moment of failure: your repeat relationship decays over the 3–5 year inter-failure gap with zero touchpoints. An appliance-age registry with lifecycle-triggered contact (maintenance tips, recall notices) keeps the thread; otherwise every failure is a fresh auction. [§ 2.3 demand management, § 4.3 lifecycle thinking, § 12 CRM]

Q10. Do manufacturer warranty-network jobs justify the rate cap, or are we subsidizing their brand promise?

A: Run the all-in: rate-capped revenue + parts margin + conversion-to-retail of warranty customers vs. fully-loaded cost. Warranty work pays only if it feeds retail conversion or fills idle capacity; as primary volume, it's their margin structure wearing your trucks. Decide by the conversion data.

[§ 11.3 coordination failure, § 4.2 utilization, § 1.3]

20. Solar installation

Q1. Sold customers refer at half the benchmark despite strong satisfaction. What happens between survey and referral ask?

A: Nothing — that's the gap. Satisfaction is passive; referral is an activity requiring a trigger, an ask, and an artifact to share. The 12–24 month post-install window (first production anniversary, first true-up bill) is when advocacy peaks, and your process has no event there. Instrument the ask. [D5, § 2.3, § 8.2 blueprint]

Q2. Why does permitting take 4× longer in our two biggest jurisdictions — and have we redesigned around it instead of absorbing it?

A: Because they're the constraint and you've treated permit lead time as weather instead of a managed process: redesign options include jurisdiction-specific design templates (pre-approved racking/plans), dedicated permit runners, and sequencing installs by permit velocity. The bottleneck governs throughput whether you manage it or not. [A3, § 7 drum-buffer-rope, § 11.2]

Q3. Are battery-attach rates limited by customer economics or sales-team comfort — which does per-rep data support?

A: Per-rep variance answers it: if attach rates vary widely by rep within the same customer economics, it's comfort and skill, not the market. Codify the high-attach reps' pitch (outage math + bill math, not technology) and retrain. [§ 6.3 stratification, § 13 standard work, § 1.1]

Q4. Why does CAC swing 2.5× between quarters — which funnel variable actually moves it?

A: Decompose CAC by channel-quarter: usually one channel (paid search or lead-buy) drives the swing as auction prices and lead quality shift seasonally, while referrals stay flat. The fix is mix management — grow the stable channels, treat volatile channels as marginal capacity. [§ 2.2 channel variance, § 13, § 1.3]

Q5. Installer production looks identical while inspection-pass rates differ 30 points between crews. What are fast crews skipping?

A: The invisible quality steps: torque verification, wire management, conduit fitting, labeling — everything the inspector checks and the production metric doesn't. Speed without pass-rate is scrap in slow motion (rework + reinspection delay). Put first-pass inspection rate on the crew scorecard beside production. [§ 6.3 SPC/Cpk thinking, § 13 Goodhart, § 6.2 internal failure]

Q6. Why do roofs flagged marginal at survey produce most of our five-year service calls?

A: Because the flag was noted and then overridden by the sale — a documented risk absorbed without pricing or redesign. The flag must become a gate: re-roof first, or price the future panel detach/reset into the contract. Unpriced flags are deferred external failure. [§ 11.4 risk assessment, § 6.2 external failure, § 0.3 specifiedBy]

Q7. When tax-credit policy shifts, which pipeline segment evaporates first — leading indicator or month-end discovery?

A: The marginal-economics segment (low usage, shaded roofs, financed deals) dies first, and the leading indicator exists: quote-stage conversion and site-survey bookings move weeks before revenue. Instrument leading funnel metrics weekly so policy shifts show up in time to react. [§ 2.2 leading indicators, § 2.1 demand uncertainty, § 11.4]

Q8. When utility rate structures change, do we re-price proposals within a week or sell old savings math until close rates force it?

A: Most installers sell the old math until the funnel screams — a governed-by-nothing pricing process. Rate schedules are public inputs; assign ownership of a rate-watch trigger that invalidates stale proposals. Selling known-stale savings math is also a trust liability waiting for a dispute. [§ 0.3 governedBy, § 2.3, § 6.1 perceived quality]

Q9. Does our financing-first presentation select for customers who default, cancel, or resent us — what does cancellation by lender show?

A: Check it: if cancellations cluster on specific lenders/dealer-fee structures, the financing frame is selecting payment-sensitive buyers who remorse-cancel at the true monthly cost. Presenting cash price first, financing second re-anchors on system value and typically halves remorse cancels. [§ 13 behavioral ops, § 2.3, § 1.1 segment selection]

Q10. Do bought leads and self-generated leads produce different system sizes, or just different acquisition costs?

A: Compare cohorts on system size, adders, financing mix, and cancel rate — not just CAC. Bought leads typically skew smaller and price-shopped; self-generated skew larger and trust-based. If confirmed, the two channels are different products and should carry different close-rate expectations and sales staffing. [§ 2.2 cohort analysis, § 1.1 MarketPositioning, § 13]

21. Dental practices

Q1. Our highest-production days coincide with highest same-day cancellation rates. What scheduling behavior creates both?

A: Overbooking to protect production: the schedule is packed past effective capacity so the day is fragile — any hiccup cascades, patients wait, and the least-committed cancel chair-side or walk. Production and cancellations are both children of the same overloaded schedule. Add buffer slots deliberately; utilization at 100% makes variability your scheduler. [A4, § 8.1 Kingman, § 7 scheduling]

Q2. Why does collections-over-production deteriorate in the same quarter every year — what payer behavior are we refusing to see?

A: Deductible seasonality: Q1 resets deductibles, patients defer or underpay, and claims adjudicate slower. It's a known seasonal demand/payment pattern you've never modeled. Plan for it: front-loaded verification, patient-pay estimates at booking, and a Q1 cash reserve. [§ 2.1 Demand Pattern seasonality, § 4.1 working capital, § 2.2 Holt-Winters]

Q3. When we bought CAD/CAM, which ROI assumption failed first: lab savings, crown volume, or same-day acceptance?

A: Usually crown volume and acceptance — the machine's economics require routing eligible cases to same-day, but scheduling templates and habit keep sending them to the lab. The technology decision lacked the infrastructural change (template redesign, presentation scripting) that realizes it. [§ 1.1 Structural vs Infrastructural decisions, § 12 DigitalOps, A10]

Q4. Treatment-plan acceptance drops when the dentist presents instead of the coordinator. Which presenter does the schedule route to?

A: The dentist, by default of authority — despite the coordinator converting better (more time, financial framing, no clinical intimidation). Routing should follow the conversion data, not hierarchy: clinical findings from the dentist, plan and financing from the coordinator. [§ 8.2 service blueprint, § 13 incentives, § 1.1 OrderWinner]

Q5. Why do insurance-directory patients leave at 18 months while referral patients stay — and does marketing know?

A: Directory patients bought access/price (qualifier shoppers); referral patients bought trust (winner-aligned). If marketing spend buys directory volume, it's acquiring churn. Shift budget toward referral-generation and measure 24-month retention by source. [§ 1.1 OrderQualifier/OrderWinner, § 2.2 cohort analysis, § 13]

Q6. Are morning huddles changing the day's outcome or an unmeasured ritual?

A: Test it: skip huddles on random days for a month and compare same-day treatment acceptance, on-time starts, and emergency absorption. If no delta, the huddle is ritual [NVA](#); if deltas appear, keep it. Sacred cows get measured like everything else. [D3 VA/NVA, § 0.3 measuredBy, § 13]

Q7. Why do associates produce 60% of owner production on the same template — clinical, communicative, or systemic?

A: Mostly communicative/systemic: associates present smaller plans, lack authority in financing conversations, and get scheduled the owner's leavings (lower-value procedures). Decompose production per hour by procedure mix first — if mixes match and production still lags, it's presentation skill, which is trainable. [§ 6.3 stratification, § 13 standard work, § 4.4]

Q8. When we add a PPO plan to grow, do new patients behave like our base — or did we import different economics?

A: You imported different economics: PPO-attracted patients are fee-sensitive, insurance-maximizing, and less loyal. That's not bad — it's a different product with different margins. The error is blending them into one P&L and one service model. Segment and price the PPO book as its own business. [§ 1.1 MarketPositioning, § 2.1 demand segmentation, § 1.3]

Q9. Hygiene reappointment looks healthy at checkout but 6-month recall show rate disagrees. Which number do we manage to?

A: The show rate — reappointment is an intention, show rate is flow. The gap is made of silent cancels and no-reminder attrition. Manage to actualized visits (the [FlowUnit](#) completing the process), and instrument the reappointment-to-show conversion as its own KPI. [§ 3.3 FlowUnit, § 0.3 measuredBy, § 13 Goodhart]

Q10. Does the unscheduled-treatment list represent clinical need or sales failure — and what do we fix first?

A: Both, separable: stratify the list by treatment urgency. Urgent-unscheduled = communication failure at diagnosis (fix the presentation); elective-unscheduled = normal patient deferral (fix the follow-up cadence). The list is demand inventory — manage its aging like any backlog. [§ 5.1 backlog as inventory, § 2.3, § 8.2 blueprint]

22. Primary care / physician practices

Q1. Why do highest-need patients generate the most after-hours calls and least revenue per visit — care model or payer mix?

A: Both, and they compound: complex patients need care management between visits (uncompensated [Activity](#)), and fee-for-service pays per visit, not per coordination. The fix is structural — CCM billing, nurse triage protocols, panel complexity weighting — not working harder inside a model that prices the wrong unit. [§ 1.1 InfrastructuralDecisions, § 8.2, D3 NNVA]

Q2. Quality-bonus metrics improve while satisfaction drifts down. Are we optimizing for the payer's definition of good?

A: Yes — Goodhart with a billing code: the payer's checklist (screenings, documentation) is measured and paid; the patient's experience (time, listening) is neither. You serve two principals with conflicting KPIs. Add experience measures to your own scorecard so the payer's don't run the practice alone. [§ 13 Goodhart, § 1.3 BalancedScorecard, § 6.3 SERVQUAL]

Q3. Why does revenue per provider vary 35% with identical panels and payer mixes?

A: Coding intensity, visit-length norms, and same-day access behavior differ — with panels equalized, the variance is in how providers work the panel. Chart-audit coding levels first (usually the biggest factor), then schedule density. This is process variation, not patient variation. [§ 6.3 SPC special cause, § 4.4 work measurement, § 0.3 measuredBy]

Q4. Does referral leakage to out-of-network specialists cost more downstream than the visit revenue protected?

A: For risk-bearing contracts, usually yes: leaked referrals lose downstream ancillary and follow-up revenue plus care coordination. Map referral flows and put a value on retention; then build the in-network preference into the referral workflow (EHR defaults, care coordinator ownership). [§ 11.2 network design, § 8.2 blueprint, § 4.1]

Q5. Did the nurse triage line cut visits by resolving issues or by adding friction — can we tell?

A: Yes: resolved issues show stable outcomes with lower visit volume; friction shows deferred visits reappearing as urgent care/ED visits later. Track downstream encounter data and complaint themes. Resolution is [VA](#); friction is [NVA](#) wearing a costume — the downstream data distinguishes them. [D3 VA/NVA, § 8.2 blueprint, § 6.3 SERVQUAL access]

Q6. No-shows cluster by appointment type rather than patient. What does the schedule do to create them?

A: Long-lead appointment types (physicals booked 6 weeks out) decay in commitment; the schedule's own lead time manufactures the no-show. Shorten booking horizons for routine types, overbook strategically by type-specific no-show probability, and confirm at 48h. [§ 8.1 abandonment, § 8.2 yield/overbooking, § 2.1]

Q7. Why does front-desk verification error rate predict denial write-offs two months later — and why isn't it managed?

A: Because verification is an upstream [Activity](#) whose defects surface downstream as denials — classic quality-at-source failure. It's unmanaged because the front desk and billing are siloed KPIs. Assign the denial cost back to the verification step and it becomes someone's number. [§ 6.3 poka-yoke/at-source, A8 traceability, § 6.2 internal failure]

Q8. Panel keeps growing while same-week access worsens. At what size did we cross the line — who was watching?

A: Nobody — panel size is an ungoverned capacity decision. The line is where appointment demand rate exceeds slot supply (Little's Law: growing panel ÷ fixed slots = growing wait). Set a panel cap per provider tied to access KPIs, and make adding patients a deliberate capacity decision. [A1 Little's Law, § 4.2 effective capacity, § 8.1]

Q9. When a provider leaves, why do 40% of patients leave the practice instead of transferring internally?

A: Because the relationship asset was personal, and your transfer process is administrative (a letter) instead of clinical (a warm handoff with the receiving provider meeting the panel). Patients treat the departure as a free choice point. Engineer the handoff like a referral, not a notice. [§ 13 relationship capital, § 8.2 blueprint, § 6.3 empathy]

Q10. Are CCM billings a care model or a billing program — would documentation survive an audit assuming the latter?

A: Stress-test it: sample CCM notes for actual care-management content (care plans, coordination calls, time logs) versus templated boilerplate. If the 20 monthly minutes aren't traceable to real activities, it's a billing program with audit exposure. Build the workflow first; the billing follows the work. [§ 0.3 specifiedBy/governedBy, A8 traceability, § 14 compliance]

23. Chiropractic

Q1. What anchors wellness-membership pricing: our cost structure or the gym membership they already pay?

A: The gym membership — customers anchor on the reference category, not your costs. Price relative to the anchor (\$40–80/mo reads as "normal wellness spend") and differentiate on clinical supervision. Cost-plus thinking loses to reference-price psychology every time. [§ 13 behavioral ops/anchoring, § 2.3 pricing, § 6.1 perceived quality]

Q2. Did adding massage therapy extend patient lifetime or repackage churn into a second line?

A: Check cohort retention: if massage clients' total relationship length is unchanged, you repackaged churn; if exit hazards drop post-massage-add, it's real extension. Cross-service attachment usually correlates with retention, but correlation here is selection — measure properly before scaling it. [§ 2.2 cohort analysis, § 13, § 1.3]

Q3. Why do maintenance patients stay for years while corrective patients with worse findings leave in months?

A: Maintenance patients bought a lifestyle identity; corrective patients bought a fix — and when pain fades (the fix working), the reason to continue vanishes. The corrective journey lacks a graduation narrative converting "fixed" into "maintained." Design the transition explicitly at the report of findings. [§ 8.2 service blueprint, § 2.3, § 13]

Q4. New-patient volume spikes after community talks while care-plan conversion falls. Who do the talks attract?

A: Free-advice seekers and the mildly curious — talks optimize for audience size, not buyer intent. Either reframe the talk (screening events with booked exams, not education seminars) or accept talks as top-of-funnel with low conversion priced in. [§ 2.3 demand shaping, § 1.1 MarketPositioning, § 13]

Q5. Why does front-desk script adherence correlate with collections more than satisfaction?

A: Because collections follow process discipline (verification, payment policy stated at booking) while satisfaction follows warmth — the script delivers the first and constrains the second. Keep the financial script, free the interpersonal moments. Measure both separately so one doesn't cannibalize the other. [§ 13 standard work + incentives, § 6.3 SERVQUAL empathy, § 0.3]

Q6. Does exam-to-report-of-findings create commitment or a decision point where we lose people — what does report-day no-show say?

A: The no-show rate answers it: high report-day no-shows mean the gap day lets urgency decay and anxiety compound. Same-visit reports (or compressed next-morning) keep the clinical momentum. The interval between diagnosis and plan is a queue with abandonment — shorten it. [§ 8.1 abandonment psychology, § 8.2 blueprint, § 2.3]

Q7. Prepaid plans complete at 80%; month-to-month patients with identical diagnoses drop by visit six. Money or mindset?

A: Both — prepaid creates sunk-cost commitment and selects for believers. You can't separate them ethically by experiment, but you don't need to: the prepaid device works regardless of mechanism. Offer it to everyone; let self-selection do the sorting. [§ 13 behavioral/sunk cost, § 2.3 pricing as commitment device]

Q8. Why do personal-injury cases dominate revenue and staff frustration — which constrains growth?

A: PI is high revenue per case but process-hostile: attorney negotiations, liens, documentation burden, 18-month payment cycles. The constraint on growth is whichever resource PI cases monopolize (usually documentation staff time). Run PI as a ring-fenced line with its own economics, not as general workload. [§ 1.2 focus, § 4.1 working capital, § 6.3]

Q9. When insurance caps visits, does outcome plateau or does documentation just stop — can we prove which?

A: Only if you keep measuring outcomes post-cap (many don't — no visit, no data). Offer a cash continuation pathway with outcome tracking; if outcomes improve post-cap, the cap was binding; if they plateau, the cap was right. Either result improves your care plans and your documentation defensibility. [§ 6.3 measurement, § 0.3 specifiedBy, § 13]

Q10. Why do reactivation campaigns work at 6–12 months gone but fail at 13–24?

A: Relationship decay crosses a threshold: within a year they still identify as your patient; past it, they've re-solved their problem elsewhere or forgotten. Reactivation is a decay function — spend the campaign budget at 6–9 months where marginal return is highest, and prevent reaching 13 at all. [§ 2.1 demand decay, § 2.3, § 13]

24. Physical therapy clinics

Q1. Why does front-desk turnover precede referral-source attrition by two quarters — is the desk the relationship holder?

A: Yes: the desk owns the daily micro-interactions with referring offices (auth updates, scheduling ease, report delivery). When the desk turns over, service friction rises silently and referrers drift. Treat front-desk tenure as a referral-channel KPI and over-invest in that seat's retention. [§ 13 relationship capital, § 8.2 blueprint/line of interaction, § 6.3 responsiveness]

Q2. Why do workers-comp cases carry worst margins and steadiest volume — what if we priced them honestly?

A: Honest pricing (true cost of auth-chasing, documentation, delays) would price you out of most comp referrals — revealing the volume is bought with margin. The strategic question: does comp volume cover fixed costs that let the profitable book thrive? If yes, cap it; if no, reprice or exit. [§ 1.3 contribution analysis, § 4.1, § 1.1 CompetitivePriority]

Q3. Switching 45→30-minute visits improved access and lowered arrival rates. What did the shorter slot signal?

A: Lower value: patients read visit length as treatment intensity, so 30 minutes feels like less care — commitment weakens, no-shows rise. Counter the signal with explicit session design ("30 focused minutes + supervised exercise"), or the access gain is eaten by the arrival loss. [§ 6.1 perceived quality, § 8.1 abandonment, § 13]

Q4. Why do plan-completing patients refer at 2× the rate of early-discharge "feeling fine" patients — does our discharge process encourage the wrong one?

A: Yes: discharging at symptom relief (not functional goal completion) produces patients who feel fine but never experienced the full transformation story worth telling. Completers become evangelists. Redefine discharge criteria around function and celebrate completion — the referral engine runs on finished arcs. [§ 6.1 perceived quality, § 2.3, § 8.2 blueprint]

Q5. Are outcomes measured to improve care or satisfy one payer — what would we measure if nobody required it?

A: That question is the diagnostic: if your honest answer differs from your current instrument set, the payer's metrics have displaced your clinical judgment (Goodhart). The ideal set — functional goals, patient-reported outcomes, visit efficiency — should exist regardless; let payer reporting be a byproduct. [§ 13 Goodhart, § 0.3 measuredBy, § 6.3]

Q6. Why does opening a second location dip the first location's volume for a full year?

A: Attention and referral leakage: the owner's presence, marketing focus, and some referrers migrate to the new site; the first location loses its differentiator (you). Plan the split as a capacity decision with dedicated demand generation for site one, or accept the cannibalization curve as the price of expansion. [§ 1.1 StructuralDecisions/Facilities, § 10 location models, § 11.4]

Q7. Why does a physician group's referral volume collapse the quarter after we miss one outcome report?

A: Because the report is the product they buy: referrers need documented outcomes to justify the referral to themselves and their patients. One miss signals unreliability in the one dimension they can't fudge. Make report delivery an automated, tracked SLA — it's an order winner, not paperwork. [§ 1.1 OrderWinner dependability, § 6.3 SERVQUAL reliability, § 12 automation]

Q8. Our best-outcomes PTs have the lowest units-per-visit. Which number does compensation reward?

A: Units-per-visit, presumably — and that's the inversion: you're paying for billing intensity while your outcomes (the referral engine) come from the opposite behavior. Redesign comp to include plan-completion and outcome measures, or your incentive system is systematically selecting against your best clinicians. [§ 13 incentives/Goodhart, § 1.3 balanced metrics, § 6.3]

Q9. Telehealth follow-ups work for some diagnoses, fail on adherence for others. Segmented or averaged away?

A: Averaged away, almost certainly. Segment by diagnosis class: exercise-adherence-dependent conditions fail remotely; education/monitoring-heavy ones succeed. Build the triage rule into the template so modality follows diagnosis, not convenience. [§ 2.1 segmentation, § 8.2 blueprint, § 12 digital ops]

Q10. Does our cancellation policy select for committed patients or flexible-job patients — what does demographic data say?

A: Check the data: strict policies filter out hourly workers and caregivers (can't guarantee schedules) — a selection effect masquerading as commitment. If your best clinical outcomes correlate with schedule flexibility, the policy is quietly rationing by class. Consider deposit-based instead of penalty-based designs. [§ 13 behavioral, § 8.2 fences, § 2.3]

25. Mental health / counseling practices

Q1. Does the group-practice model produce better outcomes or just better rent coverage — what evidence distinguishes them?

A: Outcome evidence: standardized measures (PHQ-9/GAD-7 trajectories) compared against solo-practice baselines, plus supervision/consultation effects on clinician development. If you can't produce that data, the model is rent coverage with a clinical brand — fine, but know which.

[§ 0.3 measuredBy, § 1.1, § 6.3]

Q2. Why do intake coordinators' conversion rates vary 40% with identical scripts?

A: Because scripts carry the words, not the warmth: conversion hinges on responsiveness, tone-matching, and first-available-appointment speed. Audit call recordings, find the behavioral delta (usually empathy + immediacy), and retrain on that — or centralize intake to your best converter. [§ 6.3 SERVQUAL responsiveness/empathy, § 13 standard work, § 6.3 SPC]

Q3. Are superbills/out-of-network support a growth strategy or a slow leak of payer relationships — who tracks the second?

A: Nobody tracks the second — that's the answer to watch. OON support grows per-client revenue while quietly training payers to narrow networks against you. Track payer mix shift and network renegotiation outcomes over years, not months, before calling it a strategy. [§ 11.3 coordination, § 1.1 StructuralDecisions, § 11.4]

Q4. Online bookers no-show at twice the rate of phone bookers. Friction or commitment?

A: Commitment: the phone call is a micro-relationship (voice, questions answered) that creates social obligation; the click creates none. Add commitment devices to online booking — deposits, personal confirmation calls, intake forms completed in advance — and the gap closes. [§ 13 behavioral, § 8.1 abandonment, § 8.2 blueprint]

Q5. Evening/weekend slots fill instantly and generate the highest long-term attrition. Who are we attracting?

A: Crisis-driven and convenience-shopping clients: after-hours demand includes many acute-episode starts that don't convert to ongoing therapy, plus shoppers who book multiple practices. The slots aren't bad — the expectation that they produce long-term clients is. Track cohort retention by booking slot. [§ 2.1 demand segmentation, § 8.2 yield/fences, § 13]

Q6. Why do out-of-network clients stay longer and refer more than in-network clients paying less?

A: Self-selection plus investment: OON clients chose the clinician, not the benefit, and paying more deepens commitment (sunk cost + identity). In-network clients chose "covered," and any covered therapist is interchangeable. You're selling two products; the OON product is loyalty. [§ 13 behavioral, § 1.1 OrderWinner, § 2.3]

Q7. Why does the waitlist grow while clinician utilization sits at 75% — which scheduling rule creates the gap?

A: The recurring-slot rule: standing weekly appointments lock prime slots, so utilization math includes held-but-occasionally-empty recurring hours while the waitlist can't be slotted into them. The fix is schedule hygiene: attendance policies on recurring slots and a float pool for waitlist demand. [§ 8.1 queuing/capacity allocation, § 7 scheduling, A2]

Q8. Why does marketing attract specialties we list but lack depth in — what does 3-session dropout by specialty show?

A: It shows exactly where your listing overpromises: specialties with high 3-session dropout are areas of shallow competence where clients sense the mismatch. Either build real depth (supervision, hiring) or delist — attracting demand you can't serve is negative marketing. [§ 6.3 knowledge gap, § 1.1 focus, § 2.3]

Q9. Trace the last three departed clinicians: what fraction of caseloads terminated vs. transferred — whose relationship was it?

A: Almost always the clinician's: therapy is a person-bound service product. The practice's counter is institutional trust — transfer offers with matched specialties, a warm handoff session, and brand-level continuity. If 70%+ terminate, the practice owns a building, not a patient base. [§ 13 relationship capital, § 8.2 blueprint, § 11.4 key-person risk]

Q10. Did adding psychiatry deepen therapy engagement or replace it — what does session frequency show?

A: Compare pre/post med-management session frequency for shared patients: if frequency drops after stabilization, meds partially substitute for therapy in the patient's mind; if it holds, they're complementary. Both patterns exist by diagnosis — segment before concluding, then design the integrated care pathway accordingly. [§ 2.2 cohort analysis, § 8.2 blueprint, § 6.3]

26. Optometry

Q1. Does the medical service line (dry eye, myopia mgmt) attract new patients or re-monetize existing ones — do we know?

A: Tag new-patient source by service line for six months. Medical services typically re-monetize the base first (good — higher margin, better care) and attract new patients only with deliberate external marketing. Both are valid; confusing them misallocates the marketing budget. [§ 2.3 demand management, § 1.3, § 13]

Q2. Why do contact-lens patients buy the annual supply at exam, then reorder online within the year — which does pricing acknowledge?

A: They buy once out of politeness/convenience, then price-shop the replenishment. Your pricing acknowledges the first transaction; the market owns the second. Win replenishment with subscription pricing matching online + the service layer (free exchanges, prescription monitoring) online can't offer. [§ 2.3 pricing, § 1.1 OrderWinner, § 8.2 yield/subscription]

Q3. Recall reminders work for exams, fail for contact-lens evaluations — same patient, same visit. What's different in framing?

A: The exam is health (obligation); the CL evaluation reads as a paid formality for something they already own. Reframe: the evaluation is the eye-health check for lens wearers (hypoxia, fit damage). One framing is care, the other is a toll. [§ 6.1 perceived quality, § 2.3, § 13 framing]

Q4. Remodel lifted optical sales two quarters, then baseline. Novelty, or did we stop doing something?

A: Check what changed operationally post-remodel: usually the new-display energy fades (staff stop merchandising, board goes stale). Novelty decays on its own; merchandising discipline is renewable. If sales track staff behavior, institute a board-refresh cadence as a scheduled activity.

[§ 10 layout, § 13 standard work, § 2.1 trend vs. event]

Q5. Why do PD-measured patients walk and buy online — and what's our unwritten policy on that moment?

A: Because you hand over the key to the sale (measurements + prescription) with no capture mechanism. The unwritten policy is "we pretend it doesn't happen." Written options: charge for the measurement, bundle it into lens purchase, or accept showrooming and compete on fit/adjustment service. [§ 2.3, § 6.1, § 8.2 blueprint]

Q6. Why carry 900 frame SKUs while 60 styles drive revenue — who is the board curated for?

A: For an imagined customer who values breadth; your actual customers buy the 60. Excess SKUs dilute the good ones, slow turns, and freeze working capital. Cut to a curated board (ABC by margin $\times$ turns), and depth-per-style beats breadth-of-style in conversion. [§ 5.2 ABC classification, § 4.1 working capital, § 10 layout]

Q7. Optical capture falls on days the doctor runs 15+ minutes behind. What does waiting do to willingness to buy?

A: It spends the patience budget before the purchase moment: a waiting patient enters the optical irritated, time-pressured, and emotionally done. The clinical schedule and the retail outcome are one system — protect the optical handoff (buffer slots, alert the optician to pre-engage). [§ 8.1 queue psychology, A4, § 8.2 blueprint]

Q8. Why do highest-revenue opticians have lowest frame-inventory turnover — selling skill or slow premium stock?

A: Likely both: they sell premium (slower-turning, higher-ticket) frames, so their personal mix drags turns. That's fine if intentional — but then turns shouldn't be their KPI; margin per transaction should. Match the metric to the strategy, not the reverse. [§ 13 Goodhart, § 5.2 ABC, § 1.3]

Q9. When vision plans squeeze reimbursements, we respond with volume instead of mix. What's that done to per-patient optical revenue?

A: Eroded it: volume response shortens exams, degrades the optical conversation, and pushes capture rate down — you're running faster to stay in place. The mix response (medical services, premium lenses, second pairs) raises per-patient value instead. Pull the trend; then choose. [§ 1.1 CompetitivePriority, A2 demand constraint, § 2.3]

Q10. Are second-pair sales limited by price or by no one presenting them clinically — what does attach by staff member show?

A: The staff variance answers it: wide attach-rate spread with identical pricing means presentation, not price, is the constraint. Clinical framing ("computer pair," "sun protection for your prescription") converts; discount framing doesn't. Codify the high attachers' language. [§ 6.3 stratification, § 13 standard work, § 2.3]

27. Veterinary clinics

Q1. Has anyone mapped client churn against pet age — does the 8-to-10 cliff trace to end-of-life cost shock?

A: Run the map: churn by pet-age cohort typically shows a pre-senior dip (wellness fatigue) and an end-of-life exit that never transfers to the next pet. Both are addressable: senior-care plans smooth the cost shock, and a bereavement-to-new-pet pathway retains the household. [§ 2.2 cohort/lifecycle analysis, § 2.3, § 8.2 blueprint]

Q2. Did urgent-care hours capture new demand or redistribute our own appointments into premium slots?

A: Compare total weekly visits and same-day-call volume before/after: if totals are flat while urgent slots fill with former regular bookings, you cannibalized yourself at a higher price (short-term win, long-term access damage). True capture shows new-client share rising in those hours. [§ 2.1 demand analysis, § 8.2 yield management, A2]

Q3. Why do multi-pet households generate less revenue per pet — are we discounting loyalty by accident?

A: Yes: multi-pet visits batch services (one trip, shared exam time discounts applied inconsistently) and those households price-shop harder per pet. But their lifetime value is higher — retention and share-of-pets matter more than per-pet yield. Price the household, not the pet. [§ 2.3 pricing, § 13, § 1.3 BalancedScorecard customer perspective]

Q4. Why do techs' best client-education moments happen in rooms the vet never sees — and why does comp ignore them?

A: Because comp is built on vet-visible production (services billed), while education happens in the tech's invisible labor — which drives compliance, rechecks, and retention downstream. Instrument it (education notes, compliance follow-through by tech) or your incentive system taxes your best behavior. [§ 13 Goodhart/incentives, D3 VA in service, § 6.3]

Q5. Dentals booked before checkout schedule at 40%; callback-later books at 8%. What does checkout capture?

A: Momentum and the authority moment: the vet just explained the dental disease, the pet is present, and the calendar is open. Later, the memory fades and the cost looms. Make at-checkout scheduling the default workflow — the 5× conversion gap is pure process design. [§ 8.2 blueprint, § 13 behavioral/commitment, § 2.3]

Q6. Wellness-plan clients spend more yearly but visit less for sick care. Prevention working, or access failing?

A: Test clinical outcomes: if plan pets show fewer advanced-disease presentations, prevention is working (spend shift is the design); if they're presenting later/sicker, the plan is inadvertently rationing access. The distinction determines whether the plan is care or a discount club. [§ 6.3 outcome measurement, § 2.3, § 13]

Q7. When corporate consolidators raise prices down the street, why doesn't our new-client volume move?

A: Because price isn't their customers' trigger — switching vets requires a trust rupture, not a price event, and clients don't know your price until they call. To capture their churn you must be visible at the rupture moments (moves, new pets, service failures), not just cheaper. [§ 1.1 OrderWinner trust, § 2.3, § 6.3]

Q8. Why does the reminder system fill the schedule two weeks out while same-week slots sit empty — which does capacity planning use?

A: It uses the two-week number (the full one) — and that's the misread. Same-week emptiness means demand timing mismatches slot release: hold a deliberate same-week pool for acute needs and release held slots at 48h. Managing to the two-week fill overstates your real utilization. [§ 8.2 yield management, § 7 scheduling, § 2.1]

Q9. Does the online pharmacy price-match retain clients or train them to check prices — what does refill migration show?

A: The migration data decides: if matched clients keep refilling with you, matching works; if they still drift (match is friction, autoship is one click), you're spending margin to slow an exit. The durable answer is convenience parity (autoship + home delivery), not price parity. [§ 2.3, § 8.2 blueprint, § 11.3]

Q10. Euthanasia caseload clusters on days the senior vet is off. Triage or avoidance?

A: Check who schedules them: if clients defer when told the senior vet is out, it's trust (acceptable, manage capacity); if staff quietly steer the hard appointments away from junior vets, it's avoidance — which stalls junior development and concentrates emotional load. Name the pattern, then design for it. [§ 13 behavioral, § 4.4 skill matrices, § 6.3 empathy]

28. Home health care agencies

Q1. Weekend admissions produce 3× the 30-day readmissions. What is missing from weekend intake?

A: The full clinical apparatus: weekend intake runs on skeleton staff — no supervisor review, rushed med reconciliation, no physician callback availability. The admission process's critical quality steps degrade under weekend staffing. Either staff intake to weekday standards or defer non-urgent weekend admissions. [§ 6.3 process capability, § 4.4 staffing, § 14 safety]

Q2. Why is family satisfaction high while attrition at caregiver-change events stays brutal?

A: Because satisfaction measures the relationship in steady state; the change event breaks the actual product — the specific caregiver bond. Satisfaction surveys never test the bond's portability. Reduce change frequency, do warm handoffs (overlap shifts), and measure attrition by change event specifically. [§ 13 relationship capital, § 8.2 blueprint, § 6.3 SERVQUAL]

Q3. Why do longest-tenured clients get the least care-plan review attention — what does incident rate by tenure show?

A: Complacency allocation: reviews go to new and problem clients while stable long-tenure clients quietly drift into higher acuity unassessed. Incident rate by tenure usually shows a U-shape — the tail end is your neglect showing up. Schedule reviews by calendar, not by signal. [§ 6.3 SPC stratification, § 7 planning hierarchy, § 14 safety leading indicators]

Q4. Does private-pay subsidize Medicaid or vice versa — has overhead ever been allocated honestly?

A: Run the allocation: Medicaid rates sit below true cost at most agencies once scheduling, coordination, and compliance overhead are assigned — private-pay carries it. If so, the Medicaid book is a mission/volume decision, not an economic one. Fine — but decide it as strategy, not discover it as loss. [§ 1.3 contribution analysis, § 4.1, § 1.1 CompetitivePriority]

Q5. Why do referral coordinators at our top hospital source rotate every nine months — does admission volume track their tenure?

A: Check the correlation — it almost certainly moves together. Coordinator rotation is a structural feature of hospital staffing; your response should be institutional redundancy (relationships with discharge planners, social work supervisors, not just one coordinator) so the channel survives rotation. [§ 11.4 single-point-of-failure, § 13 relationship capital, § 11.2]

Q6. Did live-in care stabilize scheduling or convert the call-out problem into a burnout problem?

A: Measure live-in caregiver tenure vs. hourly: live-in eliminates daily call-outs but concentrates strain on one person, so failure becomes total (client uncovered entirely) instead of partial. It's a reliability trade — fewer, bigger failures. Mitigate with planned relief rotations. [§ 4.3 reliability/MTBF thinking, § 4.4 job design, § 14]

Q7. Read our visit notes as an outsider: do they document staffing to the care plan or to caregiver availability?

A: If notes show visit times drifting to caregiver convenience and tasks compressing, you're staffing to availability — and an auditor or plaintiff's attorney will read it that way. The care plan is the [Specification](#); notes must trace service to spec, or the gap is liability. [§ 0.3 specifiedBy, A8 traceability, § 14 compliance]

Q8. Hospital referrals convert to admissions at half the SNF rate despite higher volume. What does the hospital handoff lack?

A: Immediacy and warmth: SNF discharges are managed (staff hand the patient to you); hospital discharges go home to a phone number and a stressed family. The gap is the first 24 hours — same-day contact and start-of-care within 48h converts; your process apparently doesn't guarantee it. [§ 8.2 blueprint, § 1.1 OrderWinner speed, § 6.3 responsiveness]

Q9. When a caregiver calls out, why does the family complaint arrive before our own missed-visit alert?

A: Because your detection is manual — someone notices a gap in a schedule grid after the fact. The family's phone call is your monitoring system, which means you have none. GPS check-in/EVV with real-time exception alerts moves detection ahead of the complaint. [§ 12 sensing layer, § 7 real-time control, § 0.1 Event]

Q10. Why does caregiver turnover concentrate in highest-acuity clients — the work or the scheduling around it?

A: Usually the scheduling: high-acuity clients get fragmented shifts, impossible drive chains, and last-minute changes because their needs are urgent — the work is meaningful, the logistics are punishing. Fix shift design around those cases (dedicated pods, guaranteed hours) before concluding it's the acuity. [§ 4.4 job design, § 13, § 7 scheduling]

29. Med spas / aesthetic clinics

Q1. Package clients use every session then rebook below pay-per-visit rates. What does the package consume besides sessions?

A: The aspiration: the package is purchased as self-promise ("I will commit to my skin"), and completing it closes the psychological project. The pay-per-visit client never closed it. Restructure packages as renewable memberships with rollover friction, or add the next-level package before completion. [§ 13 behavioral/sunk cost, § 2.3, § 8.2 subscription design]

Q2. Why do consultations convert at 70% for one provider and 35% for an identically trained other?

A: Consultation skill is not clinical training: the high converter visualizes outcomes (imaging, staged plans), prices monthly, and asks for the booking. Observe, codify, retrain — a 2× conversion spread is uncoded standard work, and it's your single biggest revenue lever. [§ 13 standard work, § 6.3 special cause, § 1.1 OrderWinner]

Q3. Would providers and marketing agree on whether clients buy results or experience?

A: Probably not — and the client buys results through the experience: they can't judge efficacy, so they judge assurance signals (environment, confidence, aftercare). Align both on: experience earns the first purchase, visible results earn the second. [§ 6.1 perceived quality, § 6.3 SERVQUAL assurance, § 1.1]

Q4. Does membership create loyalty or prepayment fatigue — what does month-13 behavior show?

A: Month-13 renewal is the verdict: high renewal = loyalty; drop-off with unused credits = fatigue. If members hold unused balances, the membership is a liability ledger, not a relationship — redesign benefits to be consumed monthly (rolling treatments) so value stays felt. [§ 2.2 cohort analysis, § 8.2 subscription, § 13]

Q5. Why does our fullest-book injector generate the lowest retail attachment?

A: Because full books punish attachment: back-to-back scheduling leaves no consultation space, and that injector optimized for throughput (which the schedule rewards). Retail attach needs a protected two minutes — redesign the template or accept the trade as deliberate. [§ 13 Goodhart, § 7 scheduling, § 2.3]

Q6. When we discount to fill the schedule, why do discounted clients occupy prime slots and full-price clients get off-hours?

A: Because discounts were offered without fences: no time restrictions, so price-sensitive clients take Saturday 11 AM while your full-price regulars get squeezed. That's yield-management malpractice — discount the off-peak inventory only, and protect prime slots for full price. [§ 8.2 revenue management/fences, § 2.3, A4]

Q7. Why do clients with complications stay at higher rates than perfect-outcome clients?

A: The recovery paradox: a well-handled complication demonstrates care under adversity — the deepest trust evidence available — while a perfect outcome is just expected. Not an argument to cause complications; an argument to engineer recovery moments (proactive check-ins, rapid response) as designed service. [§ 8.2 service recovery paradox, § 6.3 SERVQUAL responsiveness, § 13]

Q8. Best Yelp reviews name the front desk more than the injector. Who is the service, in memory?

A: The front desk — the human constants of greeting, scheduling mercy, and follow-up. The injector is assumed competent (qualifier); the desk is differentiated (winner). Staff, price, and protect accordingly: your reviews are telling you where perceived quality is manufactured. [§ 6.3 SERVQUAL, § 1.1 OrderWinner, § 8.2 line of interaction]

Q9. Why does no-show rate triple on appointments booked 30+ days out — and why do we keep booking them?

A: Commitment decays with horizon: a booking made in a motivated moment cools over a month. You keep booking them because the book looks full (a vanity metric). Counter with horizon-tiered confirmation (7-day and 48-hour reconfirmation with deposit) or shorter booking windows. [§ 8.1 abandonment psychology, § 2.1, § 13 Goodhart]

Q10. Every device follows launch spike → six-month plateau → utilization below lease payment. What if buying criteria assumed the plateau?

A: Then most devices would fail the test — which is the point. Buy for steady-state demand (plateau utilization × margin vs. lease), and treat the launch spike as marketing you could rent (demo days, manufacturer events) instead of capitalize. [§ 4.2 effective capacity/utilization, § 1.1 StructuralDecisions, § 12]

30. Urgent care clinics

Q1. Door-to-door time flat while satisfaction falls. Which interval is actually stretching?

A: Decompose the visit: usually door-to-provider improves while provider-to-disposition (labs, imaging, discharge instructions) stretches — patients feel the back half. The aggregate hides it. Instrument each interval; the patient's experience is the longest wait, not the average. [§ 3.4 process states, § 8.1, § 6.3 SERVQUAL responsiveness]

Q2. Weekday evenings carry highest acuity and thinnest staffing. Scheduling failure or market reality?

A: Scheduling failure: evening acuity is predictable (work-hours deferral), so staffing against it is a forecast error, not fate. Shift provider hours to the demand curve and add evening MAs; the pattern is stable enough to staff to. [§ 2.1 Demand Pattern, § 8.1 Erlang C staffing, § 7]

Q3. Why does provider productivity per hour vary 40% with identical patient mixes — has anyone observed it?

A: Observe first: the variance usually lives in documentation habits (charting during vs. after), rooming discipline, and MA leverage. Time-and-motion on high vs. low providers converts "personality" into trainable practice. Then set the standard. [§ 4.4 work measurement, § 13 standard work, § 6.3 special cause]

Q4. Does online check-in smooth arrivals or front-load peaks into visible rage — what does the curve show?

A: Check-in systems flatten arrival anxiety, not arrivals: if everyone checks in for 5 PM, you get a synchronized queue with informed, impatient waiters. Add capacity-shaped slotting (limited slots per interval) so the system shapes demand, not just announces it. [§ 8.1 queuing, § 8.2 yield, § 12 digital ops]

Q5. Follow-up-referral completions run below 20%. Ours clinically or the payer's contractually?

A: Clinically yours in effect regardless of contract: the patient experience breaks at the handoff. Own the loop (warm referral scheduling before discharge, 48-hour follow-up call) — completion rates double, and it's measurable differentiation for payer negotiations. [§ 8.2 blueprint, § 11.2 coordination, § 6.3 reliability]

Q6. When a competitor opens two miles away, why does volume hold but payer mix worsen?

A: They skim the commercially-insured convenience segment (better location, newer facility) while you retain Medicaid/self-pay loyalty and occupational contracts. Volume is the vanity metric; mix is the margin. Watch revenue per visit, not visits, or you'll celebrate your way into erosion. [§ 2.1 mix analysis, § 13 Goodhart, § 1.1 MarketPositioning]

Q7. Are self-pay prices calibrated to our cost per visit or the ER's billed charges — which anchor do patients use?

A: The ER anchor — patients compare to the \$1,500 ER bill they fear, not your costs. Price with the anchor in view (typically 10–20% of ER reads as miraculous value) and display the comparison; cost-plus thinking leaves the anchor's value on the table. [§ 13 anchoring, § 2.3 pricing, § 6.1 perceived quality]

Q8. When flu season surges, which breaks first: provider coverage, lab turnaround, or front desk — and do we staff to the sequence?

A: The front desk breaks first (arrival rate spikes before acuity does), then lab, then providers — but few clinics staff to that sequence. Map last year's surge hour-by-hour and pre-position the first-breaking link; surge response is a reliability calculation, not heroics. [§ 4.3 series reliability, § 8.1 Erlang C, § 2.1 seasonality]

Q9. Why do occ-med clients renew contracts while shifting injury volume to a competitor?

A: Because the renewal is administrative inertia while the volume follows the case managers' and employees' experience — you're the contract vendor, they're the preferred one. Volume drift is the truth; renewal is the paperwork. Instrument employer-level volume share, not just contract status. [§ 6.3 reliability, § 13, § 11.2 relationship vs. transaction]

Q10. Why do two-visit patients become loyalists while single-visit patients never return — what happens between visits one and two?

A: A relationship forms: the second visit means the first experience earned trust and the patient now has a "place." The lever is engineering visit two — follow-up calls, recall for unresolved issues, workplace physicals. Convert the transaction into a panel before the memory fades. [§ 8.2 blueprint, § 2.3 retention, § 13]

31. Law firms (solo/small)

Q1. When we flat-fee a matter type, why does scope creep appear in staffing notes but never invoices?

A: Because the flat fee removed the billing mechanism that made scope visible — hourly billing is also a measurement system. Flat fees require a scope spec plus a change-order trigger; without one, creep is invisible until margin audits. Define the matter's [Specification](#) and the out-of-scope price list. [§ 0.3 specifiedBy, § 9 scope control, D3]

Q2. Why track billable hours to the tenth but never referral-source conversion by matter value?

A: Because hours feed the invoice (cash) while referral analytics feed strategy (later). The measurement portfolio is all operations, no demand — and demand is the scarcer resource for a small firm. Instrument source → consult → engagement → value; it will reallocate your relationship time. [§ 0.3 measuredBy, § 13 Goodhart, § 2.2]

Q3. Why does realization fall as hourly rate rises — and where does the curve break?

A: Because rate increases push clients into scrutiny mode: higher rates trigger write-down negotiations and self-censoring on hours. The curve breaks where clients re-anchor to alternatives (typically a visible market band). Find it by plotting realization vs. rate historically; price just under the kink. [§ 2.3 pricing, § 13 anchoring, § 1.1 OrderQualifier]

Q4. Consultations ending with a clear next step convert at 2× ones with better legal analysis. What is the client buying?

A: Motion, not analysis: the client arrives anxious and buys the feeling that the problem is now in process. The next step converts the abstract ("your case is strong") into the concrete ("we file Tuesday"). Structure every consult to end with a dated action. [§ 6.3 SERVQUAL assurance/responsiveness, § 8.2 blueprint, § 13]

Q5. Does the contingency docket subsidize the hourly book or starve it of attorney hours?

A: Compute blended margin per attorney-hour across both: contingency pays in lumps years late (working-capital drag + outcome variance), hourly pays steadily. Typically contingency starves hourly of senior time while its variance goes unpriced. Ring-fence contingency capacity with explicit risk-adjusted return targets. [§ 4.1 working capital, § 1.3, § 11.4 risk]

Q6. When a paralegal leaves, why does matter profitability drop for two quarters under the replacement?

A: Because the paralegal carried the matter's working memory — unwritten status, client quirks, document locations. The replacement bills learning time to the same matters. Codify matter state (checklists, status dashboards) so the knowledge lives in the system, not the person. [§ 13 standard work/knowledge, § 4.4, § 11.4 key-person risk]

Q7. Avvo clients churn after one matter; CPA-referred clients return for years. Does intake scoring know?

A: Make it know: tag source at intake and score expected relationship length into acceptance and fee decisions. Avvo delivers transaction-shoppers; CPA referrals arrive embedded in an advisory relationship. Same service, two demand classes — treat them differently in pricing and follow-up. [§ 2.1 demand segmentation, § 2.2 cohort, § 1.1]

Q8. What stretches between signing and document delivery — and how many estate clients does the gap lose?

A: The draft-review ping-pong: unsigned drafts wait on client input with no deadline, and your queue discipline lets paying-now work jump ahead of waiting-on-client work. Set aging SLAs on draft states and automate client nudges; the gap is where referrals die silently. [§ 3.4 Blocked/Starved states, § 7 dispatch rules, A1 Little's Law]

Q9. Why does the partner's personal book grow while firm-originated matters shrink — succession risk or comfort?

A: Both, feeding each other: the partner takes the best inbound (comfort), so the firm's brand never develops its own rainmakers (risk). Measure origination share deliberately; if the firm can't originate without the partner, you have a job, not a firm. Build associate origination into comp. [§ 13 incentives, § 1.1 StructuralDecisions, § 11.4 succession risk]

Q10. Our best-written engagement letters correlate with worst collections. Who reads them that closely — and pays slowly?

A: Sophisticated, litigation-minded clients: they read closely because they manage vendors adversarially, and they pay slowly because the letter gives them the terms to exploit. The letter isn't causing slow pay — it's selecting for clients who negotiate everything. Watch that cohort's realization, not their compliments. [§ 2.1 segmentation, § 13 behavioral, § 0.3 governedBy]

32. Accounting / bookkeeping / tax prep

Q1. Why do cleanup quotes overrun 60% on price-shopping clients?

A: Adverse selection: price-shoppers have the messiest books (they churned through providers or DIY'd) and the least tolerance for revised estimates. Add a paid diagnostic phase before quoting cleanup — the diagnostic prices the uncertainty that the flat quote currently eats. [§ 2.1 uncertainty, § 9 risk, § 13 adverse selection]

Q2. Are tax-planning deliverables read — and if we stopped the PDF, who would notice?

A: Test cheaply: skip it for a willing cohort and track inquiries. If few notice, the deliverable is a ritual — replace the PDF with a 15-minute review call, which clients experience as value. Deliverables nobody reads are [NVA](#) dressed as service. [D3 NVA, § 6.1 perceived quality, § 8.2 blueprint]

Q3. When the client's bookkeeper is our main contact, why does renewal depend on their tenure more than our quality?

A: Because the bookkeeper is your translator and protector — they chose you, defend you, and interpret your work upward. When they leave, the new person brings their own vendor preferences. Build relationships above and beside the contact; one-thread accounts are rented, not owned. [§ 13 relationship capital, § 11.4 key-person risk, § 11.2]

Q4. Why does effective hourly rate on fixed-fee work decline yearly despite price increases?

A: Because scope grows faster than fees: each year clients add "one more report," and your team absorbs it. Fixed fees need annual scope re-specification, not just price escalation — re-baseline the [Specification](#) every renewal, or the fee buys more matter every year. [§ 0.3 specifiedBy, § 9 scope creep, § 1.3]

Q5. Advisory upsells attach to mid-size clients, never smallest or largest. What does mid-size have?

A: Pain plus budget plus a decision-maker you can reach: small clients can't pay, large clients have in-house staff. Mid-size feels the operational pain and owns the checkbook — the sweet spot. Stop pitching advisory outside it; deepen inside it. [§ 1.1 MarketPositioning, § 2.1 segmentation, § 2.3]

Q6. Staff error rates rise in October, our "quiet" month. What is October actually full of?

A: Extension deadlines and year-end planning prep — invisible workload. The calendar says quiet; the work mix says compressed deadlines on complex returns. Map workload hours (not return count) by week; October is a second tax season you're not staffing for. [§ 2.1 demand pattern, § 7 aggregate planning, § 4.4 workload]

Q7. Does CAS pricing reflect the software stack cost or the client's freed time — which story closes?

A: The freed-time story closes; the stack-cost story justifies. Price on the client's alternative (hiring a bookkeeper at 3× your fee), and present software as included infrastructure. Value-anchored pricing wins; cost-anchored invites negotiation. [§ 2.3 pricing, § 13 anchoring, § 6.1]

Q8. Why do we lose clients to the nephew with QuickBooks and win them back 18 months later — what's never mentioned in the return conversation?

A: The mess you cleaned up silently: the nephew's errors, the penalties avoided, the categories fixed. Your competence is invisible because it prevents visible events. Document the rescue in the win-back — "here's what we found and fixed" converts the return into a retention lesson. [§ 6.1 perceived quality of prevention, § 2.3, § 13]

Q9. Why does monthly-client churn concentrate in months 10–14, right after the first full tax cycle?

A: Because the tax cycle is the annual value audit: the client sees the total year's fees at once and re-evaluates. If the year's touchpoints were silent (books done invisibly), the audit fails. Engineer quarterly visible value (review calls, savings found) so the audit has evidence. [§ 6.1 perceived quality, § 8.2 blueprint, § 2.3]

Q10. Extension clients are our most profitable relationships, yet we treat extensions as failures. Which belief is wrong?

A: The failure belief: extensions are a product — deadline-flexible clients who pay premium attention rates and generate second-season work. Treating them as defects corrupts your scheduling (you cram them into April anyway) and your client language. Build the extension season as a deliberate second calendar. [§ 1.1 emergent strategy, § 7 capacity planning, § 13]

33. Marketing / advertising agencies

Q1. Why do referrals arrive exactly 13 months after engagements end — what does the lag say about value visibility?

A: Your value becomes legible only in the client's year-over-year comparison after you're gone — during the engagement it's obscured by noise and attribution fog. Compress the lag: instrument results into the client's own dashboards during the engagement so value is attributed while you're still there to be renewed. [§ 6.1 perceived quality, § 0.3 measuredBy, § 13]

Q2. When the client hires internally, why does our contract survive if we trained them and die if we didn't?

A: Because training repositions you from labor to force multiplier: the trained hire needs your strategy and escalation, not your hands. Untrained, the hire replaces you. Make client-enablement a designed part of the engagement — it converts your biggest threat into your retention mechanism. [§ 1.1 VerticalIntegration/make-vs-buy from client's view, § 13 knowledge, § 8.2]

Q3. Why does utilization run hot while project margins run cold — and which clients invert it?

A: High utilization on underpriced scope: you're busy doing unbilled work — overruns absorbed as "client service." The inverting clients (hot utilization AND margin) have tight scopes and change-order discipline. Margin leaks live in the gap between sold scope and delivered scope; meter it weekly. [§ 9 EVM-style scope control, § 1.3, D3]

Q4. When we report great metrics and the client cancels anyway, what were they buying that we stopped delivering?

A: Confidence in the relationship's future — metric reports are backward-looking; clients renew on belief about next quarter (strategy, ideas, attention). The report proves the past; the cancellation mourns the absent future. Lead reviews with what's next, not what happened. [§ 6.3 assurance, § 1.1 OrderWinner, § 13 Goodhart on reporting]

Q5. Clients signed in hot streaks churn faster than slow-quarter signings. What does the hot-streak pitch promise?

A: Whatever it takes: in hot streaks, sales optimism outruns delivery scoping — you sell the case-study outcome as the expected outcome. Slow-quarter pitches are sober and set survivable expectations. Audit win-rate vs. expectation-inflation; your delivery team is paying sales' hot-streak debt. [§ 13 incentives, § 6.3 knowledge gap, § 11.3 overreaction]

Q6. Why do strategists' best ideas appear in lost pitches more than signed accounts?

A: Because lost pitches are unconstrained by the client's budget and politics — the ideas were too big for the deal. Two fixes: price-tier the strategy so ambition is purchasable, and mine the lost-pitch archive for your signed clients' next proposals. [§ 1.1 MarketPositioning, § 2.3, § 13]

Q7. Does blended-rate pricing hide which services lose money — which line would we kill if we knew?

A: Almost certainly yes: blended rates cross-subsidize (usually production subsidizes strategy, or vice versa). Run rate-card costing per service line for a quarter; agencies that do this typically find one line earning under loaded cost — and killing or repricing it is pure margin. [§ 1.3 costing, § 4.1, § 0.3 measuredBy]

Q8. Scope creep concentrates in highest-paying retainers — absorbed from exactly the clients who could afford change orders. What are we protecting?

A: The relationship fantasy: fear that charging your best clients invites re-evaluation. Inverted logic — large clients respect metered scope (their own procurement works that way). Absorption teaches them scope is free, which trains the creep. Change-order your best clients first. [§ 11.3 incentive misalignment, § 9 scope control, § 13]

Q9. Why do case studies come from clients whose results we can least attribute?

A: Because big-name clients have visible results (they'd grow anyway) and logos that sell; your attributable wins are with unglamorous clients. This is survivorship marketing — fine for credibility, dangerous if your own team starts believing the attribution. Know the difference internally. [§ 13 selection bias, § 6.1 perceived quality, § 1.1]

Q10. Read our SOWs as a meter: do they bill for judgment or production volume?

A: Production, typically — deliverable counts, hours, posts per month. Judgment is what clients actually renew for, but it's unbillable as scoped. Shift SOWs toward outcomes and decision-support (strategy days, testing roadmaps) so the meter measures what retains. [§ 0.3 measuredBy/specifiedBy, § 1.1 OrderWinner, § 2.3]

34. IT services / MSPs

Q1. Why does tools-stack cost per tech rise as we add clients — which integration tax are we not counting?

A: The per-client onboarding and exception-handling tax: every non-standard client environment creates one-off scripts, alert tuning, and documentation in your RMM/PSA — tools are priced per endpoint but operated per exception. Enforce stack standardization or price the exceptions; the tax is real either way. [§ 12 platform layer, § 13 standard work, § 1.3]

Q2. Why do we win on security posture, then get reviewed quarterly on response time?

A: Because the pitch set one winner and the QBR measures another: post-sale, the client's daily experience is tickets, not posture. Align the quarterly review to include the security outcomes you sold (incidents prevented, patch compliance) — or the measured metric becomes the de facto product. [§ 13 Goodhart, § 1.1 OrderWinner migration, § 6.3]

Q3. Why does churn concentrate among clients who never had a major incident, while disaster-rescued clients stay forever?

A: Invisible value: no incident means no proof of your prevention — the quiet client concludes "nothing ever happens, why pay?" The rescued client saw the alternative. Counter with value reporting (blocked threats, patched CVEs, downtime avoided) — make prevention legible or it reads as absence. [§ 6.1 perceived quality of prevention, § 8.2 recovery paradox, § 2.3]

Q4. Does per-seat pricing survive client headcount downturns — or is every layoff an unbudgeted renegotiation?

A: It doesn't: per-seat couples your revenue to their HR decisions — a demand risk you've priced as zero. Options: minimum-commit tiers, per-device + floor pricing, or explicit annual true-up terms. Choose the structure before the downturn chooses the renegotiation for you. [§ 11.4 demand risk, § 2.3 pricing structure, § 0.3 governedBy]

Q5. When the client's office manager is our champion, why does the account wobble when they go on leave?

A: Single-thread dependency: your institutional knowledge of the client lives in one person who routes everything. When they're out, requests stall, tickets age, and the client experiences your service as degraded. Multi-thread every account — at least two contacts who know the escalation paths. [§ 13 relationship capital, § 11.4 single-point failure, § 8.2 blueprint]

Q6. Are QBRs changing client behavior or billing ourselves for theater — which clients act on them?

A: Audit: which clients implemented last quarter's recommendations? If <20%, QBRs are theater for most accounts. Restructure: only run QBRs where an executive owner attends and decisions get minuted; for the rest, a quarterly report suffices. Match the ceremony to the audience's authority.

[D3 NVA, § 6.3, § 1.1 segment]

Q7. Clients resisting stack standardization become our highest-margin accounts. What does resistance predict?

A: Deep entanglement: resistant clients have complex environments where you're embedded (custom integrations, institutional knowledge) — high switching costs and premium tolerance. Resistance predicts lock-in. Profitable, but note the risk mirror: you're also locked into their complexity. [§ 11.2 switching costs, § 11.4, § 1.3]

Q8. Why does project revenue spike right before managed-service churn — are projects a symptom of disengagement?

A: Yes, often: a client planning their exit commissions final projects (migration to in-house, new platform) before canceling the agreement. The project pipeline is a leading churn indicator nobody watches. Flag accounts with unusual project activity for a relationship review. [§ 2.2 leading indicators, § 13, § 11.2]

Q9. Fully-managed clients generate more tickets per seat than break-fix converts. Engagement or noise — which are we staffed for?

A: Mostly engagement (they use the service they pay for) plus expectation inflation (everything is "included"). Staffing assumes break-fix ticket rates and drowns in managed volume. Normalize: ticket deflection (self-service, automation) and per-seat ticket budgets baked into pricing tiers. [§ 8.1 arrival rates, § 12 automation, § 2.3]

Q10. Escalations cluster by time of day, not complexity. Who is on shift — or what do we call complex?

A: Both possibilities resolve with data: if escalations cluster at shift boundaries or junior-coverage windows, it's coverage skill, not ticket nature. If truly time-correlated regardless of staff, it's client-side patterns (their batch jobs, their EOD panic). Stratify by shift first — the cheaper fix. [§ 6.3 stratification, § 7 scheduling, § 4.4 skill matrices]

35. Management consulting

Q1. Why do weekly-steering engagements overrun more than monthly ones — what does cadence signal?

A: Weekly cadence signals client-side unreadiness: frequent meetings are demanded where decisions are contested and sponsorship weak — the meetings multiply alignment work instead of progress. Cadence is a diagnostic, not a preference. Price weekly-cadence engagements for the political work they actually contain. [§ 9 project management, § 13, § 1.1]

Q2. Would clients describe what they bought the way we describe what we sold — answers or confidence to act?

A: Confidence to act: the deck is the artifact, the purchase was permission and cover for a decision. This matters for delivery: a "correct" answer the client can't act on is a failed output, regardless of analytical quality. Design deliverables for actionability (owner, first step, political path). [§ 6.1 perceived quality, § 0.3 Output spec, § 13]

Q3. Why do smallest engagements produce the largest follow-ons — and why doesn't the pipeline model reflect it?

A: Because small engagements are paid diagnostics of trust: the client buys a low-risk sample, and the follow-on is the real purchase. Your pipeline weights proposals by size, undervaluing the seed engagements. Model expected lifetime value, not first-contract value, in pursuit decisions. [§ 2.2 leading indicators, § 1.3 BalancedScorecard customer perspective, § 13]

Q4. Proposals win when they name the political problem, lose when they name only the operational one. Which are we hired to solve?

A: Both, but the political one is the order winner: the operational problem justifies the budget, the political problem decides the vendor. Discovery must map the decision politics (who loses if this works?) or your beautiful operational answer dies in committee. [§ 1.1 OrderWinner, § 13 behavioral, § 6.3 knowledge gap]

Q5. Junior utilization looks great while their work gets rewritten. Who pays for the rewrite?

A: The firm does, invisibly: senior rewrite hours hit the engagement as unbilled time, so margins absorb what utilization metrics hide. Measure rework loops explicitly (draft → rewrite cycles per deliverable) — high junior utilization with high rewrite is negative productivity wearing a good KPI. [§ 13 Goodhart, § 6.2 internal failure/rework, § 1.3]

Q6. Why does thought-leadership generate inbound in industries we don't serve and silence in those we do?

A: Because your content is generic-smart — it resonates where nobody can evaluate it and says nothing your target industry doesn't already know. Targeted insight requires industry-specific claims (risky, research-heavy). Generic content is cheap to produce and worthless for positioning; pick a vertical and be specifically right. [§ 1.1 MarketPositioning/focus, § 13, § 2.3]

Q7. Why do validate-a-decision engagements produce best references and worst margins?

A: Best references because the client got cover for what they wanted (you're the hero); worst margins because scope is political and endless — every stakeholder needs a meeting, a revision, a caveat. Price validation work for the consensus-building it actually is, not the analysis it pretends to be. [§ 2.3 pricing, § 9 scope, § 13]

Q8. Repeat business collapses when the sponsoring executive changes, despite documented results. Whose asset were the results?

A: The sponsor's — results entered the client's mythology as the executive's win, and your relationship was personal. Institutionalize during the engagement: multiple stakeholder ownership, results embedded in the client's operating reviews, successors briefed before transitions. [§ 13 relationship capital, § 11.4 key-person risk, § 8.2]

Q9. When a deliverable gets shelved, was the analysis wrong or the client unable to act — do our diagnostics separate them?

A: Build the separation in: end every engagement with an implementation-readiness assessment (decision authority, budget line, owner named). A shelved deliverable with no readiness assessment tells you nothing; with one, shelving becomes predictable — and preventable in scoping. [§ 0.3 Output/Specification, § 9 handoff, § 6.3]

Q10. Does day-rate pricing select for advice-valuing clients or budget-spending clients — which renews?

A: Budget-spenders buy days to consume a line item and vanish at year-end; advice-valuers buy outcomes and renew. Day rates make you interchangeable with the budget cycle. Shift pricing toward value/retainer structures that select the advice-valuers — they're the annuity. [§ 2.3 pricing structure, § 1.1 segment selection, § 13]

36. Staffing / recruiting agencies

Q1. Our fastest fills produce the highest 90-day fall-offs. What does speed cost the match?

A: Verification depth: fast fills skip the second-interview probe, reference triangulation, and expectation alignment — the activities that predict retention. Speed is a qualifier clients demand; match quality is the winner they remember at day 91. Sell the slate with a speed-quality frontier made explicit. [§ 1.1 OrderQualifier/Winner trade-off, § 6.3, § 13]

Q2. Why do exclusive searches fill slower than contested ones in our own data?

A: Urgency asymmetry: contested searches create competition-driven pace (you fear losing the fee); exclusives relax into thoroughness — and client responsiveness also drops without rival pressure. The pitch for exclusivity is quality; the data says manage exclusive timelines contractually or drift is structural. [§ 13 incentives, § 11.3 coordination, § 7 scheduling]

Q3. Follow one req through the ATS: is the workflow optimized for the candidate or the client?

A: For the client — status fields, submission formats, and client-portal updates dominate; candidate communication is where reqs go silent. Yet candidate experience is your supply chain's supplier relations: ghosted candidates become unavailable talent. Instrument candidate-touch SLAs into the workflow. [§ 11.2 supplier management analog, § 8.2 blueprint, § 12]

Q4. Our best clients pay late; our worst pay on time. What does payment data reveal?

A: That your best clients treat you as a strategic partner with negotiated terms (they pay late because they can and you allow it — relationship depth), while transactional clients process you through AP automation. Payment behavior is a proxy for relationship type — price the late payers' terms into the fee. [§ 4.1 working capital, § 11.2 contract terms, § 2.1 segmentation]

Q5. Does temp-to-perm build loyalty or train clients to treat us as a free trial — what does conversion pricing show?

A: If conversion fees are waived or discounted to close deals, you're running a free-trial program — clients cycle temps perpetually instead of converting. Conversion pricing is the product's fence: weaken it and temp-to-perm becomes churn-by-design. Hold the fee; discount only for volume commitments. [§ 8.2 fences, § 2.3 pricing, § 13]

Q6. Client concentration grows despite a diversification goal. Which account behavior rewards the drift?

A: The big account's own gravity: more reqs, faster feedback, dedicated attention feels earned — every operational incentive pulls toward the whale. Diversification fails because it's a goal with no mechanism. Set a concentration ceiling with a pricing premium above it, or the drift is the strategy.

[§ 11.4 concentration risk, § 13 incentives, § 1.1 deliberate strategy]

Q7. Ghosting rates vary more by recruiter than by market. What do low-ghosting recruiters do?

A: They manage the in-between: pre-close candidates on the offer specifics, maintain contact during notice periods, and prep counteroffer conversations. Ghosting is a process failure at the offer-to-start gap. Codify the low-ghosters' touchpoint cadence into the playbook. [§ 8.1

abandonment, § 13 standard work, § 6.3 special cause]

Q8. Why do high-submission recruiters have the lowest sendout-to-placement ratios?

A: Because submission volume is the visible activity metric, so they spray resumes and let the client screen — offloading your quality function onto the client's time. The ratio exposes it.

Rebalance KPIs toward sendout-to-placement and fall-off rates, or the ATS fills with noise. [§ 13 Goodhart, § 0.3 measuredBy, § 6.3 quality at source]

Q9. When a client demands 10 resumes and hires from the first three, what are we being paid for?

A: Confidence through contrast — the client needs to see the field to trust the choice. The seven extra resumes are decision-theater [NVA](#), but refusing them reads as withholding. Educate:

calibrated slates of 3–4 with market data attached sell better than volume. [D3 NVA, § 13 behavioral, § 6.3 assurance]

Q10. When unemployment falls, why does our gross margin fall instead of rising with demand?

A: Because scarcity shifts power to candidates and clients simultaneously: clients resist fees ("you have no candidates"), candidates extract concessions, and your cost per placement (search hours) rises. Tight markets reward retained/exclusive models and punish contingency — reposition before the cycle, not during. [§ 2.1 demand cycles, § 11.2 contract structure, § 1.1]

37. Web design / development shops

Q1. Does our tech stack serve the client or our hiring pipeline — could the client tell?

A: Check your stack choices against client needs: if you're building marketing sites in the framework your developers want on their resumes, the stack serves hiring. Clients can't tell today; they discover it at handoff when nobody local can maintain it. Match stack to client

maintainability. [§ 1.1 ProcessTechnology decision, § 13 principal-agent, § 0.3 specifiedBy]

Q2. Why do redesign clients take their next redesign to a different agency — what happened at handoff?

A: The relationship ended at launch: no maintenance contact, no performance reviews, no roadmap conversation — three years of silence, then a fresh RFP. The next redesign is seeded the day this one launches. Own the post-launch lifecycle or forfeit the next cycle. [§ 2.3 retention, § 8.2 blueprint, § 11.2]

Q3. Discovery predicts project success better than portfolios, yet we discount it first. What signal are we selling cheap?

A: The certainty signal: discovery is where the client learns you understand their problem — the actual purchase decision happens there. Discounting it says "the thinking is free, the typing costs money," inverting your value. Price discovery as a product (it de-risks both sides) and conversion improves. [§ 1.1 OrderWinner, § 2.3 pricing, § 6.1]

Q4. When the client's team can't update their site six months post-launch, whose failure contractually and commercially?

A: Contractually theirs (you delivered what was specced); commercially yours (their frustration finds your name at every future RFP). Training and admin UX are **NNVA** only until the client calls them broken. Bake handoff training and a 90-day support window into the spec. [D3 NNVA, § 0.3 specifiedBy, § 8.2 blueprint]

Q5. Fixed-price projects overrun by almost exactly 30% regardless of size. Which scope variable drives it?

A: Content: client-provided copy, images, and approvals arrive late and wrong-sized in every project, independent of site size. The 30% constant is the content-and-feedback loop you never bound. Contract for content deadlines with schedule consequences, or price the 30% in as the content-risk premium. [§ 9 critical chain/buffers, § 11.2 client dependencies, § 13]

Q6. When we lose maintenance clients, is it to cheaper providers or their new in-house hire — which should change pricing?

A: Diagnose per exit: losses to in-house hires are structural (they outgrew the product) — respond with escalation-tier services; losses to cheaper providers mean your maintenance tier is over-served (cut scope/price). Two different churns need two different responses; exit-interview the difference. [§ 2.1 churn segmentation, § 1.1, § 2.3]

Q7. Why quote hosting/maintenance as afterthoughts when they carry highest margin and lowest cost-to-serve?

A: Because sales focuses on the visible project (the site) while the annuity (hosting, care plans) is where the business model actually lives. Recurring revenue is a different product with a different sale — attach it at proposal as a designed tier, not a checkbox. [§ 1.1 StructuralDecisions, § 2.3, § 8.2 subscription]

Q8. Our best developer generates our worst client satisfaction. Does staffing route them correctly?

A: No: technical excellence and client-facing patience are different resources, and your routing treats them as one. Put the best dev on architecture and internal escalation, shield them from client calls, and staff client contact with communicators. Role-fit beats hierarchy. [§ 4.4 job design/skill matrices, § 6.3 SERVQUAL, D5]

Q9. Clients who approve designs fastest request the most revisions in development. What is fast approval telling us?

A: That they didn't evaluate: fast approval means the client deferred judgment until the design became real (clickable, in-browser). The revision flood is the evaluation happening late. Add a structured design-walkthrough gate that forces decisions before sign-off. [§ 9 stage gates, § 13 behavioral, § 6.3 acceptance criteria]

Q10. Are accessibility/performance sold as risk mitigation or features — which framing survives the budget meeting?

A: Risk mitigation: features get cut when budgets tighten; lawsuit exposure and SEO/conversion penalties survive. Frame accessibility as legal-risk reduction and performance as revenue protection (conversion loss per second of load time), and both stop being optional line items. [§ 2.3 framing, § 11.4 risk, § 6.1]

38. SEO / digital agencies

Q1. When a Google update hits, why do strong-content clients suffer the same volatility as weak ones — what's differentiation worth?

A: Because volatility is platform risk, not quality risk: both clients ride the same algorithm. Your differentiation shows in recovery speed and floor protection, not immunity. Sell resilience (diversified traffic, owned channels) instead of implying updates won't hurt — honesty here is retention. [§ 11.4 platform/disruption risk, § 6.1 perceived quality, § 2.1]

Q2. If we had to stake our fee on one number — rankings, traffic, or revenue — which would we choose?

A: The honest answer is the one with controllable causation: rankings you can move but they're decoupled from value; revenue is value but multi-causal. The question exposes the reporting game: you report what's favorable, not what you'd bet on. Close that gap or a client's CFO eventually will. [§ 13 Goodhart, § 0.3 measuredBy, § 1.3]

Q3. Why do we sell 12-month commitments but staff accounts as if they might leave monthly?

A: Because staffing to contract length is risky and staffing light is safe — but the client experiences under-staffing as under-delivery and fulfills your fear by leaving. The staffing IS the retention strategy: front-load senior attention in months 1–4 when judgments form. [§ 13 self-fulfilling incentives, § 1.1 InfrastructuralDecisions, § 2.3]

Q4. Does our content produce pipeline or just traffic — have we ever seen the client's CRM?

A: Get CRM access or concede the question: traffic-to-pipeline attribution requires the client's conversion data. Agencies that never ask are choosing comfortable deniability. Make pipeline visibility a contract term — it protects you (proof) more than it exposes you. [§ 0.3 measuredBy, § 11.2 information sharing/CPFR analog, § 1.3]

Q5. Why does sales promise what delivery calls worst-case scope — where does translation break?

A: At the handoff artifact: sales sells outcomes in language, delivery inherits a contract in deliverables. The gap lives between the pitch deck and the SOW. Require delivery sign-off on proposals above a threshold, and make the SOW quote the pitch's promises verbatim. [§ 0.3 specifiedBy, § 11.3 incentive misalignment, § 9 handoff]

Q6. Clients rank for promised keywords and cancel citing no business impact. What did we sell that rankings were to deliver?

A: Revenue, implicitly — you let the client believe rankings → traffic → pipeline without owning the last two conversions. The keyword was achievable; the expectation was the product defect. Sell the chain explicitly (with conversion-rate assumptions and landing-page work) or sell rankings as visibility only, priced accordingly. [§ 6.3 knowledge gap, § 1.1 OrderWinner, § 2.3]

Q7. When a client's team asks to learn from us — future referral source or future cancellation?

A: Depends on what you teach: teach tactics and you train your replacement; teach strategy and you become their permanent advisor. History says: enable the client's team and you move up the value stack; hoard and you're a cost line awaiting an in-house hire. [§ 13 knowledge strategy, § 11.2, § 1.1 VerticalIntegration from client side]

Q8. Reporting cadence increases before cancellations. Retention tool or anxiety generator?

A: Anxiety generator: the pattern is usually the agency's defensive response to wobbly accounts — more reports, more meetings, more proof. The client reads activity as nervousness. If the account is wobbling, the fix is a strategy conversation, not reporting volume. [§ 13 Goodhart, § 6.3 assurance, § 2.2 leading indicator: this is one]

Q9. Why do link deliverables show clean in reports while traffic moves with brand searches we didn't cause?

A: Because brand demand (PR, word of mouth, the client's product) drives the searches that produce your "organic growth," while your link work moves metrics nobody buys on. Separate branded vs. non-branded organic in every report — it's the difference between your effect and their momentum. [§ 13 attribution honesty, § 0.3 measuredBy, § 6.3]

Q10. Local-service clients churn at months 7–9 regardless of results. What expectation does month-one onboarding install?

A: The six-month miracle: onboarding that promises "results in 6 months" schedules the client's evaluation at month 7 — right when local SEO is compounding but not yet visible. Reset the expectation curve at onboarding (realistic timeline, leading indicators by month 3) or the calendar itself fires you. [§ 13 expectation management, § 6.3 knowledge gap, § 2.2 leading indicators]

39. HR / payroll services

Q1. When clients ignore our advice and terminate anyway, why does renewal probability rise?

A: Because the termination goes sideways as predicted, and your warning becomes proof of value in retrospect — you're the one who knew. Documented advice is an option that pays at the next crisis. Keep advising and keep the paper trail; renewal happens at the vindication. [§ 6.3 assurance, § 11.4 risk advisory value, § 13]

Q2. Which mental account does the client's CFO file us under — insurance premium or purchased service?

A: Insurance premium, most likely: paid monthly, invisible when things work, resented at renewal. Insurance framing caps your price and invites shopping. Shift the frame to purchased outcomes (audits survived, hours saved, claims avoided) with an annual value report. [§ 13 framing/mental accounting, § 6.1 perceived quality, § 2.3]

Q3. Why do policy deliverables get bought by clients who never implement them — are we selling documents or cover?

A: Cover: the handbook is purchased as lawsuit armor ("we have policies"), not as an operating tool. If you want to sell documents, price accordingly; if you want impact (and referrals), sell implementation — training, acknowledgment tracking, annual review. Know which product you're in. [§ 6.1 perceived vs actual quality, § 0.3 Output spec, § 2.3]

Q4. Why is NPS high while share-of-wallet stays flat for years?

A: Satisfaction without expansion: clients like the service they bought and have no idea what else you sell. NPS measures the relationship with the current product; wallet share needs cross-sell mechanics (QBRs with expansion agendas). High NPS is permission, not a plan. [§ 13 Goodhart, § 2.3 demand expansion, § 1.3]

Q5. Clients forward our compliance alerts to their lawyers. Which of us is trusted for what?

A: You're trusted for detection (fast, practical), the lawyer for validation (authority, liability). That's a fine division — unless the lawyer starts offering the detection. Position as the operational layer (alerts, implementation, training) and ally with employment counsel rather than competing on authority. [§ 1.1 MarketPositioning, § 11.2 complementors, § 6.3 assurance]

Q6. Does the PEO-comparison business cannibalize consulting — and which do our referrals point to?

A: Check referral flow: if referrals arrive for "PEO shopping" and leave as PEO placements, your consulting brand is feeding the brokerage. Cannibalization is fine if brokerage margins beat consulting and you choose it — dangerous if it's accidental. Decide which business the brand serves. [§ 1.1 deliberate vs emergent strategy, § 1.3, § 13]

Q7. Under-15-employee clients churn on price; over-50 churn on service. What happens in the danger zone between?

A: The worst of both: 15–50 clients outgrow your cheapest tier's service but balk at premium pricing — they're forming HR sophistication without the budget to match. Design a specific mid-tier product (dedicated rep lite, quarterly reviews) or the zone is a churn funnel. [§ 2.1 segmentation, § 1.1 ProductProcessMatrix analog, § 2.3]

Q8. When a client's HR manager quits, why do we inherit the job without the budget?

A: Because you're the continuity default — the client assumes the vendor absorbs the gap. Without a contract mechanism, absorbed work is free work. Add an interim-HR services rider (priced, time-boxed) triggered by HR-staff vacancies; the inherited job is a product you're not selling. [§ 0.3 governedBy, § 11.4 client-side disruption, § 2.3]

Q9. Payroll accuracy is 99.9% while satisfaction tracks response time. What are clients scoring?

A: The experience of exceptions: accuracy at 99.9% means the relationship lives in the 0.1% — and what clients remember is how fast a human answered when their payroll broke. Accuracy is the qualifier; responsiveness is the winner. Staff the exception desk like it's the product. [§ 1.1 OrderQualifier/Winner, § 6.3 SERVQUAL responsiveness, § 8.2 recovery]

Q10. Crisis-onboarded clients are our stickiest; steady-state buyers churn. What does panic create that prudence does not?

A: Vivid value: the crisis client watched you absorb their chaos (filing back-taxes, fixing misclassification) — a demonstrated rescue. The steady buyer never saw you do anything. Create crisis-equivalent visibility for quiet accounts: annual "what we prevented" reviews and compliance audits that surface saved risk. [§ 8.2 recovery paradox, § 6.1 perceived quality of prevention, § 13]

40. Business coaching

Q1. Does our framework create results or dependency — what do clients do in month 13 if we stop?

A: That's the test: clients who keep running the operating rhythm after you leave bought capability; clients who collapse bought accountability-as-a-service. Both are legitimate products — but price dependency as a subscription (it ends when you do) and capability as a build (it transfers). Confusing them creates refund-shaped resentments. [§ 13 knowledge transfer, § 1.1, § 2.3]

Q2. Why does our own business violate the systems we teach — which promise would fail a client audit?

A: Usually the metrics cadence and the delegation ladder — coaches sell operating discipline they don't run. The audit risk is reputational: one client visiting your operation collapses the authority. Run your own system minimally but genuinely; authenticity is the product's core quality dimension. [§ 6.1 perceived quality/assurance, § 13, A10 consistency]

Q3. Why do accountant referrals convert to engagements while coach referrals convert to coffee?

A: Referral-source credibility transfers: the accountant's referral arrives with financial pain documented and trust attached; the coach's referral is a favor between peers with no urgency. Invest in referrer types who see the pain in the numbers (accountants, attorneys, bankers). [§ 2.2 channel quality, § 6.3 assurance transfer, § 11.2]

Q4. Best case studies come from industries we've never marketed to. What does inbound know that positioning doesn't?

A: That your method fits a buyer you haven't named: case-study industries share a structure (owner-operated, operational pain, cash to spend) your positioning never articulated. Mine the inbound pattern, then position deliberately toward it — the market already voted. [§ 1.1 MarketPositioning, § 2.2 pattern analysis, § 13]

Q5. When the client's spouse attends a session, why does retention double — and why isn't that in intake?

A: Because the spouse converts private homework into household commitment — the business owner's decisions survive at the dinner table. Retention doubles because the change effort gains a stakeholder. Make spouse/partner inclusion a designed step in onboarding, not an accident. [§ 13 behavioral/commitment, § 8.2 blueprint, § 2.3]

Q6. If accountability came as software at a tenth of our fee, what would clients say they miss?

A: The pattern recognition and the being-seen: software can remind, but it can't say "this is the third quarter you've deferred the hiring decision, and here's what that costs." Your durable value is judgment + relationship, not reminders. Price and sell the judgment; let software carry the reminders. [§ 1.1 OrderWinner, § 12 automation boundary, § 6.3]

Q7. Why does the roster skew toward declining businesses when marketing targets growth-stage owners?

A: Because declining owners feel pain and growth owners feel fine — pain buys, comfort browses. Either embrace the turnaround market (position for it honestly) or reframe growth marketing around growth's pains (scaling chaos, founder bottleneck), which are real but need naming. [§ 2.3 demand reality, § 1.1 positioning, § 13]

Q8. Clients who implement least attend most. Which behavior does pricing reward?

A: Attendance: flat monthly fees meter presence, not progress, so the talkative client who never implements is your best-paying customer. Consider milestone-linked structures (phase gates tied to implementation) — you keep revenue aligned with the value you claim to sell. [§ 13 Goodhart, § 2.3 pricing structure, § 0.3 measuredBy]

Q9. When revenue jumps, clients credit the market and cancel; when it stalls, they blame themselves and stay. Which attribution does renewal depend on?

A: Self-blame — meaning your renewal economics currently depend on clients feeling inadequate, a fragile and slightly toxic basis. Counter with counterfactual attribution during wins: instrument what changed because of the work (decisions made, systems installed) so success has your fingerprints on it. [§ 13 attribution, § 6.1 perceived quality, § 0.3 measuredBy]

Q10. Why do group-program clients out-satisfy 1:1 clients at a third of the price — what does that do to the premium tier?

A: Peer effects outperform advisor attention for most owners: the cohort provides accountability, normalization, and ideas at scale. It means 1:1 must be re-scoped to what only 1:1 does (confidential, complex, political work) — or the premium tier is a legacy product riding brand inertia. [§ 1.1 ProductProcessMatrix, § 2.3 tiering, § 13 community]

41. Restaurants (independent)**Q1. Would we want to know if our most loyal customers' habits lose us money?**

A: Yes — run it: loyalists who camp for hours on a single entrée consume table-time (your true capacity unit) at low yield. But before acting, price the whole relationship (frequency × lifetime + referral value). If they still lose money, the fix is gentle yield design (table-time norms at peak), not eviction. [§ 8.2 yield management, § 1.3 BalancedScorecard customer perspective, § 4.2 capacity]

Q2. Does third-party delivery cannibalize direct orders or reach new customers — what does address overlap show?

A: The overlap data is the answer: high overlap = cannibalization at 25–30% commission (you're paying platforms to serve your own customers); low overlap = genuine reach. If cannibalizing, steer direct (owned ordering incentives, pickup pricing) and fence delivery to acquisition zones. [§ 2.1 channel analysis, § 8.2 fences, § 11.3 double marginalization]

Q3. 86'd dishes correlate with specific prep-shift assignments, not demand spikes. What is prep not doing?

A: Pars: the prep sheet's quantities aren't being executed or aren't computed from actual usage — certain prep cooks eyeball instead of weigh/count. The 86 map is the audit: items 86'd on specific shifts trace to specific pars failures. Enforce measured pars with sign-off. [§ 5.1 cycle stock/pars, § 13 standard work, A5 conservation of flow]

Q4. Friday revenue grows while Friday profit shrinks. Which line item moves opposite sales?

A: Labor and waste: Friday volume triggers overtime, extra prep (over-production for the rush), and expediting chaos — the incremental Friday dollar carries higher marginal cost than the average dollar. Compute marginal Friday profitability; the answer may be fewer covers, better sequenced. [§ 1.3 marginal analysis, A4, § 7 labor scheduling]

Q5. Why do highest-rated items carry lowest contribution margins — has anyone put the lists side by side?

A: Do it — that's menu engineering: plot popularity × contribution. Your stars are probably low-margin (signature dish priced for fame); your profit hides in puzzles nobody orders. The fix is classic: reprice, re-engineer plate cost, or reposition high-margin items into the spotlight. [§ 1.3 contribution analysis, § 2.3 menu as demand shaping, § 13]

Q6. When the chef is off, food cost improves but reviews fall. What is the chef actually spending on?

A: Quality insurance: better ingredients, wider margins of safety on portions, refuse-to-serve substandard plates. The sous-chef runs the spec tighter and serves the borderline plate. Neither is "right" — but now you can price the chef's quality premium and decide if the review protection is worth the food-cost points. [§ 6.2 CoQ trade-off, § 6.1 quality dimensions, § 13]

Q7. Why do private-event inquiries convert at 20% while walk-ins convert at 100% — what does our phone manner optimize?

A: The phone manner optimizes for ending the call: no availability calendar at hand, no package prices ready, follow-up promised "by email tomorrow." Event inquiries are your highest-value demand and get your lowest-quality process. Build the inquiry script (date, headcount, budget, instant proposal) and conversion doubles. [§ 8.2 blueprint, § 1.1 OrderWinner responsiveness, § 2.3]

Q8. Why do highest-tipped servers sell the least wine — which behavior does training reward?

A: Speed and warmth (which drive tips) over suggestive selling (which risks friction). Tips and check-size are different behaviors; your training rewards the visible one. If wine margin matters, decouple: train pairings as hospitality (not sales) and measure attachment separately from tips. [§ 13 Goodhart, § 2.3 upsell design, § 6.3]

Q9. Regulars increasingly order delivery instead of dining in. What did the dining room stop giving them?

A: Occasion value: when the room stopped feeling like a reward (service routine, recognition faded, ambiance dated), the food alone competes with their couch — and loses on convenience. The dining room must sell what delivery can't: theater, recognition, pace. Audit the regulars' in-room experience. [§ 6.1 perceived quality/experience, § 8.2 blueprint, § 2.3]

Q10. Why does labor cost spike on slow Tuesdays, not busy Saturdays?

A: Fixed skeleton crew against near-zero revenue: Tuesday runs the minimum staff (kitchen lead + one server) on 30% of Saturday's sales — the labor percentage explodes on the slow day, not the busy one. Schedule to demand curves per daypart; slow days need demand-shaping (specials, events), not just resignation. [§ 7 labor scheduling, § 2.1 daypart patterns, § 2.3]

42. Coffee shops & cafes

Q1. Pastry waste runs 30% while the same items stock out weekly. What is par-setting optimizing?

A: Nothing — pars are static while demand is day-of-week and weather volatile. Waste and stockouts coexisting is the signature of no forecast model: set pars by day-of-week demand history (newsvendor logic for perishables), and bake to a two-batch schedule (open + mid-morning). [§ 5.2 Newsvendor, § 2.1 demand pattern, A4]

Q2. When a chain opens nearby, why do we lose the afternoon crowd first, not the morning one?

A: Morning is habit-locked (commute ritual, low involvement); afternoon is discretionary (meetings, treats) and comparison-shopped. Your defensible core is the morning ritual — deepen it (speed, recognition, subscriptions); the afternoon needs a distinct reason to exist (workspace terms, event programming). [§ 13 habit/behavioral, § 1.1 OrderWinner, § 2.3]

Q3. When we raise drink prices, transaction count holds but average ticket falls. What are customers trading down from?

A: The add-ons: pastry attachment and size upgrades die first under price pressure — the drink is the habit, the extras are the elastic margin. Counter with bundle pricing (drink + pastry at a combined anchor) so the trade-down path leads to a designed option, not just less. [§ 2.3 price architecture, § 13 anchoring, § 1.3]

Q4. Morning regulars visit daily but generate less margin than the weekend crowd we can't retain. Which customer does the menu serve?

A: Neither deliberately — the menu evolved. Decide: if regulars are the annuity, optimize for speed and loyalty pricing; if weekend margin matters, build the weekend occasion (brunch items, family space). Serving both with one menu and one flow means each subsidizes the other's experience. [§ 1.1 focus, § 8.2 yield, § 2.1 segmentation]

Q5. Does the laptop policy change revenue per seat-hour — have we ever measured seat-hour revenue?

A: Measure it: seat-hour revenue is your true productivity metric (tables are the capacity resource). Laptop campers typically halve it at peak. Options: time-limited tables at peak, laptop-free zones, or minimums. You can't choose without the per-seat-hour number. [§ 8.2 yield management, § 4.2 capacity, § 1.3]

Q6. Why do mobile orders raise volume but degrade the in-store metrics that built the brand?

A: Because mobile demand jumps the same production line: baristas serve the queue screen instead of the counter guest — one bottleneck, two arrival streams, and the in-person customer watches drinks go to invisible people. Separate the flows (dedicated mobile station/pickup shelf) or throttle mobile slots at peak. [A3 bottleneck, § 8.1 two-class queues, § 8.2 blueprint]

Q7. Why do fastest-ticket baristas have lowest greeting scores — and which predicts repeat visits?

A: Speed and warmth trade off at the individual level, but repeat visits track recognition (name, usual order) more than raw speed — up to the wait threshold. Ideal: fast enough to respect the commute, warm enough to be a ritual. Staff the morning rush for throughput and train the 10-second recognition habit. [§ 6.3 SERVQUAL responsiveness vs empathy, § 13, § 2.3]

Q8. Seasonal drinks outsell the core for two weeks, then underperform. Are we buying inventory for the spike or the plateau?

A: For the spike, if purchasing is driven by launch excitement — and the plateau then becomes waste. Buy for the plateau with spike coverage via flexible suppliers or smaller initial lots; treat the launch curve as a newsvendor problem with a known decay. [§ 5.2 Newsvendor, § 2.1 demand lifecycle, § 13 recency bias]

Q9. Read the kitchen labor budget as a confession: coffee business with food, or food business with coffee?

A: The budget confesses the truth your branding may deny: if food labor and COGS approach 40%+, you're a cafe-restaurant with food-level complexity and coffee-level ticket averages. Either grow average ticket to match the kitchen's ambition or simplify the kitchen to match the ticket.

[§ 1.1 emergent strategy, § 1.3, § 10 layout/process choice]

Q10. Why is loyalty redemption so low it functions as an unpaid discount — and what happens when friction drops?

A: Low redemption means the program is breakage-based: it works financially only because rewards expire unused. Dropping friction raises redemption and visit frequency — test whether incremental visits cover the redemption cost. If yes, friction was costing you the actual point (frequency); if no, you had a quiet margin. [§ 13 behavioral, § 2.3 loyalty design, § 1.3]

43. Bars & taverns

Q1. Why does the kitchen close early on highest-traffic nights — what does the 10 PM crowd buy instead?

A: Nothing from you — they buy elsewhere or just drink, and the drunk-hungry demand (your highest-margin, least-price-sensitive food demand) walks to the pizza place. The kitchen closes on labor-cost logic on exactly the nights when marginal food revenue is highest. Extend food hours on peak nights; the data will embarrass the old rule. [§ 2.3 demand accommodation, § 1.3 marginal analysis, § 7]

Q2. Happy-hour customers stay for dinner at neighboring restaurants, not ours. What does our room fail to offer at 7 PM?

A: A transition: lighting, music, and menu that say "dinner now." Your room stays a bar; the neighbor's becomes a restaurant. Engineer the 6:30–7 PM state change (menu drop, lighting shift, table service) so the same customer experiences the next occasion without leaving. [§ 8.2 service blueprint/states, § 6.1 aesthetics, § 2.3]

Q3. Why do inventory counts reconcile in slow months and never in busy ones?

A: Because variance hides in volume: busy months generate over-pouring, comps, spillage, and theft that exceed counting tolerance, and nobody investigates when sales look good. Count high-velocity SKUs weekly (perpetual inventory on the top 20) — slow-month accuracy is false comfort. [§ 5.1 inventory accuracy, § 13 Goodhart, § 6.3 control limits]

Q4. Does the cocktail program justify its labor — or do three drinks fund the menu?

A: Run contribution per drink: elaborate programs typically show 3–5 cocktails carrying the category while the craft tail consumes prep labor and spoilage. Keep the signature performers, batch-prep what you can, and cut the romantic tail unless it's genuinely the brand. [§ 1.3 contribution, § 5.1 perishables, § 1.1 focus]

Q5. When a bartender quits, why does their shift revenue drop 15% for months — was the margin ever ours?

A: Partly no: relationship-driven pours (their regulars, their heavy hand) leave with them. The structural fix: institutionalize regulars (house loyalty program, multiple bartenders knowing names) and standardize pours (jiggers, POS-pour tracking) so margin belongs to the house. [§ 13 relationship capital, § 6.3 standardization, § 5.1]

Q6. Why do we lose customers to the newer place every spring and win them back by fall?

A: Novelty decay cycle: the new venue offers change; you offer reliability. The annual churn-and-return is a market rhythm — instead of fighting it, schedule your own novelty events (menu refreshes, pop-ups) into their spring window to shorten the defection. [§ 2.1 cyclical patterns, § 13 novelty, § 2.3 counter-programming]

Q7. Why do regulars' tabs subsidize promo nights attracting one-time crowds — which customer does the calendar serve?

A: The promo calendar serves vanity metrics (busy-looking nights). Compute per-event incremental profit: most promos discount existing demand (regulars would've come) while attracting deal-seekers who never return. Design promos for off-nights with fences that protect regular pricing. [§ 8.2 fences, § 13 Goodhart, § 2.3]

Q8. Busiest nights produce worst pour-cost percentages. Theft, over-pour, or chaos pricing — which scales with volume?

A: Over-pour and chaos: under rush, free-pouring drifts heavy and POS ringing lags (drinks served unringed, then comped to reconcile). Theft exists but doesn't scale cleanly with volume. Test: measured-pour pilots on peak nights and ring-before-pour discipline — the variance will confess. [§ 6.3 special cause, § 13, § 5.1 variance analysis]

Q9. When we book live music, revenue rises but per-customer spend falls. Which pays the band?

A: Neither automatically — compute the event P&L: cover charges + incremental bar margin vs. band cost. Music fills the room with lingerers (low per-hour spend) and displaces some high-spend diners. It pays when it activates dead nights or drives covers; it taxes you when it crowds out spending. [§ 8.2 yield, § 1.3, § 2.3]

Q10. Are trivia/karaoke nights building a crowd or renting one — what does the following Wednesday look like?

A: The following Wednesday is the test: if baseline Wednesdays improved after trivia launched, you're building; if trivia crowds vanish on non-trivia nights, you're renting (the crowd belongs to the format, not your bar). Renting is fine if the rent is profitable — just know which. [§ 2.1 demand analysis, § 13, § 1.3]

44. Catering companies

Q1. Why do headcounts come in 10% under guarantee while food cost comes in over?

A: Because you prep to the guarantee plus a safety factor, and the guarantee is inflated by clients protecting themselves — you pay for both buffers. Contract for a final-count deadline with a prep-to number, and price the guarantee window explicitly. [§ 5.1 buffer/safety stock, § 0.3 governedBy contract, A4]

Q2. Why do staffing agencies send our own former employees at 1.5× their old wage — and why do we pay it?

A: Because event staffing is a peak-capacity problem you solve weekly with panic instead of a bench: the agency premium is the price of your missing labor buffer. Build an on-call bench (alumni list, cross-trained part-timers) — the 1.5× is financing your workforce flexibility at payday-lender rates. [A4 capacity buffer, § 4.4 labor flexibility, § 7 aggregate planning]

Q3. Our largest events carry the thinnest margins. Where does event size stop paying?

A: Where coordination complexity outruns scale economies: big events add rental logistics, staffing agency mixes, and client-committee revisions — complexity grows superlinearly while per-head price stays flat. Find the kink in your historical size-vs-margin curve and price beyond it with a complexity surcharge. [§ 1.2 product-process fit, § 9 project complexity, § 1.3]

Q4. Are drop-off orders a gateway to full-service or a replacement — what do second events show?

A: The second event tells you: clients whose second event upgrades = gateway; clients repeating drop-off = you've trained them that drop-off suffices. Design the upgrade path (drop-off clients get staffed-service trials at cost) or the gateway becomes a ceiling. [§ 2.1 cohort analysis, § 2.3 product ladder, § 13]

Q5. Our highest-rated events generate the fewest referrals. What does the referral-producing event look like?

A: One with a visible host: corporate events (the organizer's reputation rides on you, and they tell peers) and weddings (guests experience you live). Your highest-rated events may be high-end but private. Engineer referral surfaces: branded touches guests notice, follow-up with the organizer's network. [§ 2.3 referral mechanics, § 6.1 perceived quality visibility, § 13]

Q6. Why do tastings convert corporate at 80% and weddings at 40% — what is the wedding client deciding?

A: The wedding client is deciding between dreams, not vendors — the tasting is one input to an emotional, committee-driven (partner, parents, Pinterest) decision. Corporate decides on competence, which the tasting demonstrates. For weddings, sell the planner relationship and the visual story, not just the food. [§ 1.1 OrderWinner by segment, § 6.3 assurance, § 13]

Q7. Does venue-preferred status drive bookings or cap pricing — what do non-preferred bookings at those venues pay?

A: Compare: if preferred bookings price below your direct bookings at the same venues, the status is a volume-for-margin trade (kickbacks, rate expectations). Preferred status is a channel — manage it like one: know its toll and its conversion, and keep a direct channel alive beside it. [§ 11.2 channel economics, § 2.3, § 11.3]

Q8. Corporate rebooks annually; weddings are one-time by definition. What does our sales effort assume about the mix?

A: If sales effort chases weddings (glamorous, big tickets) while corporate quietly compounds, the effort is backwards: corporate accounts are annuities worth years of margin. Balance acquisition cost accordingly — a corporate account justifies 5–10× the pursuit budget of a wedding. [§ 1.3 BalancedScorecard customer perspective, § 2.3, § 1.1]

Q9. When a client cuts budget mid-planning, which lines do we protect — and is that policy written?

A: It isn't, and it should be: protect the guest-visible experience (food quality, service staffing) and sacrifice the invisible (rental upgrades, garnish complexity) — degradation the client chooses should never become quality the guests blame on you. Write the downgrade ladder into the contract process. [§ 0.3 governedBy, § 6.1 perceived quality, § 9 change control]

Q10. Why is kitchen utilization below 50% while we decline weekend work for lack of staff?

A: Two different constraints misread as one: kitchen capacity is idle midweek; weekend failure is labor, not kitchen. You're kitchen-rich and crew-poor. Sell weekday production (corporate lunch programs, meal prep, commissary rental) to monetize the kitchen, and fix weekend staffing separately. [A3 different bottlenecks, § 4.2 utilization, § 2.3 counter-cyclical demand]

45. Food trucks

Q1. Why do slowest weekdays generate the best catering leads — and why cut that day first?

A: Because slow days mean visible, unhurried presence — event planners and office managers actually stop and talk. The day's value is demand generation, not sales. Price the day as marketing (cost per catering lead) before cutting it; it may be your cheapest acquisition channel. [§ 2.3, § 1.3 opportunity-cost accounting, § 13]

Q2. If a buyer valued this business tomorrow, which asset would they pay for: a brand that can go brick-and-mortar, or a route that can't?

A: The brand: routes die when you stop driving; brand demand (followers, catering reputation, product recognition) transfers and scales. Build the transferable assets deliberately — recipes documented, social accounts on business infrastructure, catering contracts assignable. [§ 1.1 StructuralDecisions, § 13 knowledge codification, § 4.1 intangible resources]

Q3. Our most profitable locations have the shortest lines. Throughput or margin — which are we optimizing?

A: Margin, apparently — long lines can mean underpriced demand (queue as rationing) while short-line spots charge full price to committed buyers. Neither is wrong, but know: lines are marketing and volume; short lines are yield. Choose per location rather than averaging the confusion. [§ 8.1 queue as signal, § 8.2 yield, § 10 location]

Q4. When the truck breaks down, why don't followers migrate to catering — branding failure or channel failure?

A: Channel failure: the follower relationship lives in the location ritual ("the truck at Fifth and Main Tuesdays"), not in a contactable asset. You never captured the audience (email/SMS list), so downtime makes you invisible. Build the owned channel while the truck runs. [§ 11.4 disruption resilience, § 12 platform/CRM, § 2.3]

Q5. Why do organizers negotiate our fee down while their application process selects for highest-priced trucks?

A: Because selection and negotiation are separate games: they select for quality (draw), then squeeze margin (their revenue share logic). Your counter is exclusivity and data — bring your per-event sales history and negotiate from proven draw, or pay for placement only where the crowd converts. [§ 11.3 coordination/bargaining, § 2.3, § 13]

Q6. Two popular items share one fryer — a bottleneck never priced into location choices. What does it cost per stop?

A: Compute it: fryer-constrained throughput $\times$ contribution per item $\times$ queue-abandonment rate at that location. If line length caps sales at busy stops, the fryer is your true constraint and either menu engineering (shift mix to grill items at high-volume stops) or a second fryer pays for itself. [A3 bottleneck, § 3.4 Blocked state, § 10]

Q7. Does the social following predict location turnout, or just which posts people like — what does geo-data say?

A: Check follower geography vs. stop locations: followings are usually metro-wide while stops are hyper-local — a post reaches fans 40 minutes away who'll never come. Geo-targeted announcements (per-stop SMS lists, location stories) convert; broadcast likes don't. Measure turnout lift per channel per stop. [§ 2.2 causal indicators, § 13 vanity metrics, § 10]

Q8. Why do commissary costs rise faster than route revenue — which stops pay their share?

A: Commissary is fixed-cost demand on a variable-revenue route: low-revenue stops consume the same prep, storage, and drive economics as good ones. Allocate commissary + drive cost per stop; the bottom tier of stops is subsidized by the top. Cut or reprice the bottom. [§ 10 routing/location economics, § 1.3, § 4.1]

Q9. Why do festivals beat brewery routes per day but not per month — which calendar do we plan against?

A: Per month is the truth: festivals are rare, fee-heavy, weather-exposed spikes; breweries are the recurring base. Plan capacity against the recurring base and treat festivals as margin opportunities with explicit go/no-go economics (fee $\div$ expected margin). [§ 2.1 demand patterns, § 1.3, § 11.4 weather risk]

Q10. Why does a second truck raise total revenue 40% instead of 100% — where did the 60% go?

A: Into shared constraints: your prep capacity, your management attention, and cannibalized stops (second truck splits existing territories). The first truck ran on founder oversight; the second runs on systems you don't have. Build the commissary and ops layer before truck three. [A3 shifting bottleneck, § 1.1 StructuralDecisions, § 13]

46. Bakeries

Q1. Are seasonal items profit or tradition — do they displace rack space from higher-margin staples?

A: Compute rack-space contribution: seasonal items earn spike revenue but displace staple capacity and add SKU complexity (molds, ingredients, training). If seasonal margin-per-oven-hour beats staples, keep; many bakeries find the opposite and keep items for identity — fine, but call it marketing and budget it as such. [§ 1.3, § 4.2 capacity allocation, § 6.1 identity value]

Q2. Why does the Saturday line coexist with barely-break-even weekly profit?

A: Because the line is Saturday-only: peak-day queues mask five thin days. Weekly profit = demand distribution problem, not a Saturday success story. Options: shift demand (weekday subscriptions, office delivery), or cut midweek capacity (hours, staff) to the true curve. [§ 2.1 demand distribution, § 8.2 yield, § 7 capacity]

Q3. As wholesale grows, the retail display empties by noon. Which business are we in?

A: A capacity-constrained one choosing between channels by default: wholesale (steady, lower margin, volume) is eating production that retail (higher margin, brand-building) needs. Set channel allocations deliberately — reserve retail display stock, or admit wholesale is the business and shrink the storefront promise. [§ 1.2 DecouplingPoint/channel allocation, A2, § 1.1 focus]

Q4. Does the bread program exist for profit or identity — what if artisan labor were priced honestly?

A: Usually identity: fully-loaded artisan bread (4 AM labor, slow fermentation space, waste) rarely beats pastry margins. Identity is a legitimate reason — bread is the brand's proof — but then judge it as marketing spend with brand metrics (traffic, halo on other sales), not as a product line. [§ 1.3, § 6.1 perceived quality/identity, § 13]

Q5. Why do cake tastings book solid while orders concentrate in two months — what happens to the other ten months of tasters?

A: They're price-and-date shopping: tasting demand is year-round (engagements are), but weddings cluster — most tasters choose other dates, other vendors, or smaller cakes. Convert off-peak tasters with off-peak pricing and non-wedding occasions (anniversaries, milestone birthdays) to fill the valleys. [§ 2.1 seasonality, § 2.3 counter-seasonal demand, § 8.2 yield]

Q6. Why does the day-old discount move product but train customers to wait — can we see both effects?

A: Yes, with timing data: if day-old buyers used to buy fresh (same households shifting), the discount cannibalizes; if they're a distinct price-sensitive segment, it's pure capture. Fence it: day-old available only after a set hour and never on the fresh shelf beside full-price. [§ 8.2 fences, § 13 anchoring, § 2.1 segment]

Q7. When ingredient costs spike, why do we shrink portions instead of raising prices — which do customers notice?

A: They notice shrinkage more, eventually — the regular's croissant got smaller reads as betrayal, while a 25-cent increase reads as inflation (external, forgivable). Behavioral research is clear: shrinkflation damages trust more than honest pricing. Raise the price; keep the product. [§ 13 behavioral/shrinkflation, § 6.1 perceived quality, § 11.3]

Q8. Custom cakes carry 60% margins and consume 80% of owner attention. What's the actual hourly return?

A: Compute it honestly: margin minus the owner-hours at a real wage. Custom work is usually below shop-average hourly return once consultation, sketching, and revision time are counted — it survives on love. Either price consultations and revisions, or accept it as the craft showcase and cap its volume. [§ 1.3 true costing, § 4.4, § 2.3]

Q9. Why does decorators' best work appear in photos while repeat orders cluster on simpler designs?

A: Because the portfolio optimizes for wow (acquisition) while repeat buyers optimize for taste, price, and reliability (retention). Both are real demand — but staff the photo work as marketing (limited slots) and systematize the simple designs (templates, junior decorators) as the volume engine. [§ 2.1 segmentation, § 1.2 product-process fit, § 13]

Q10. Pre-order customers show up late or not at all at rates walk-in data never predicted. What does prepayment do to obligation?

A: Full prepayment without reminders reduces felt obligation for some (paid = settled, mentally closed), while deposits keep the loop open. Either way: the fix is operational — reminder messages the day before with pickup windows and a release policy. Measure no-show by payment type, then design. [§ 13 behavioral, § 8.1 abandonment, § 0.3 governedBy policy]

47. QSR franchise units

Q1. Why does crew turnover cluster at month 3 — what does exit data say we get wrong in week 2?

A: Week 2 is when training ends and reality starts: the new hire hits full-speed stations during rushes with thinning support, and the job's actual demands (pace, heat, rude customers) arrive before any belonging does. Fix the week-2 experience: graduated station exposure, a named buddy, and a 30-day check-in. [§ 4.4 job design/onboarding, A9 learning curve, § 13]

Q2. Which decisions does the franchise agreement actually leave us — when does the operator become a tenant of the brand?

A: You're left with labor scheduling, local marketing, maintenance timing, and wage levels — the margin levers, not the revenue levers. The tenant line is crossed when brand mandates (remodels, LTOs, hours) consume your remaining autonomy's return. Audit decisions-you-own vs. profit-you-own; that ratio is your real position. [§ 1.1 StructuralDecisions scope, § 11.3 coordination, § 13]

Q3. Why does the remodel's sales lift last four months while the debt lasts ten years?

A: Because the lift is novelty + relaunch marketing, not structural demand — the market renormalizes while the amortization doesn't. The true test: comp-store sales vs. control at month 12. If flat, the remodel was a brand-tax, not an investment; negotiate scope or timing accordingly. [§ 2.1 novelty vs trend, § 1.3 capital ROI, § 13 recency]

Q4. Does the aggregator mix help the P&L — or does franchisor reporting hide the commission line?

A: Pull store-level aggregator P&L directly (commission + packaging + promo funding + incremental labor) vs. in-store margin. Franchisor dashboards often report delivery as gross sales. If delivery orders are 20%+ of volume at negative contribution, you're subsidizing the platform with your food cost. [§ 11.3 platform economics, § 1.3 contribution, § 13 Goodhart]

Q5. Why do two stores three miles apart cannibalize catering but not lunch traffic?

A: Lunch is geographic-habit demand (location-bound); catering is a considered purchase where both stores quote the same buyer — you're bidding against yourself. Assign catering territories or centralize catering sales so the stores stop competing on the only demand that travels. [§ 10 location/coverage models, § 2.1 demand type, § 11.2]

Q6. Late-night loses money while corporate mandates hours. What would opting out cost — has anyone priced it?

A: Price both sides: late-night fully-loaded cost (labor + utilities + shrink at 15% utilization) vs. the franchise penalty or brand-standards consequence. Often the penalty is negotiable or smaller than the bleed. Present corporate with the unit economics — franchisors respond to data from their top operators. [§ 1.3, § 4.2 utilization, § 11.2 contract negotiation]

Q7. Unit economics beat disclosure averages while morale lags them. Is one causing the other?

A: Likely: above-average economics in QSR are frequently bought with below-average labor spend (thin shifts, no slack), which the morale metrics record. The risk is the lagging indicator — turnover and service failures arrive after the margin was banked. Watch if your economic edge survives fully-loaded turnover cost. [§ 13 leading vs lagging, § 4.4, § 6.2 internal failure]

Q8. Best shift leads promoted to manager quit within a year at the role's salary cap. What is the promotion for?

A: Currently: removing your best shift leader and creating a salaried burnout. The role's economics (flat salary + unlimited hours + modest cap) punish exactly the conscientious. Redesign: real authority, bonus on store metrics, and a ceiling conversation before promotion, not at resignation. [§ 4.4 job design, § 13 incentives, A9]

Q9. When corporate launches a promotion, why does food cost rise more than traffic justifies — do we have data to push back?

A: Build it: per-promo unit economics (promo-item food cost %, waste from LTO ingredients, cannibalization of full-price items). Franchisees who present per-store promo P&Ls get heard; those who complain get ignored. Your POS has the data — the analysis is the missing activity. [§ 1.3, § 13 Goodhart (corporate's metric is system sales), § 2.1 cannibalization]

Q10. Drive-thru speed improves during mystery-shop weeks, then decays. What is the measurement training?

A: Theater: the team performs for the measurement window, meaning the capability exists but isn't sustained — your daily management system (not the mystery shop) sets real behavior. Replace episodic inspection with continuous timer data reviewed daily; what gets watched daily stays fast. [§ 13 Goodhart, § 6.3 SPC continuous control, § 7 real-time control]

48. Pizza shops

Q1. When a chain runs 50%-off week, why do we lose price-sensitive customers permanently, not temporarily?

A: Because the promo week breaks their habit and the chain's app captures them (saved order, rewards points) — the trial converts to a new default. Counter before promos: lock your price-sensitive base into your own loyalty habit so a discount week is just a week. [§ 13 habit formation, § 2.3 loyalty design, § 11.4 competitive shock]

Q2. Friday ticket times degrade 12 minutes while staffing peaks. Which station is the bottleneck?

A: Almost always the oven: staffing adds labor, but oven capacity is fixed — Friday demand saturates bake time while cut-and-box backs up behind it. The fix is oven-side: par-bake strategy, oven management discipline, or channel throttling (delivery promise times). More staff feeding a full oven just crowds it. [A3 bottleneck, § 3.4 Blocked, § 10 line balancing]

Q3. Our best-selling location has our worst delivery economics. Is the radius drawn around demand or history?

A: History: radii are inherited from opening day. Delivery economics = drive time per drop × drops per run; dense zones close-in beat far sprawl orders. Re-draw by marginal delivery cost per zone (drop the tail, tighten promises in the core) — the top-line dip is usually tiny against the cost recovery. [§ 10 location/routing, § 1.3, § 4.1]

Q4. When we source better cheese, reviews improve but repeat rates stay flat. What are repeat customers loyal to?

A: The total ritual: price point, speed, the familiar taste profile — quality upgrades they notice in the moment don't change the reorder decision, which runs on habit and value perception. The better cheese was worth it for differentiation at acquisition, not retention. Know which you're buying. [§ 6.1 perceived quality dimensions, § 13 habit, § 2.3]

Q5. Why do lunch-slice regulars never convert to dinner delivery — are they the same market?

A: No: lunch-slice buyers are workers solving speed and price; dinner delivery is household occasion demand. Different [FlowUnit](#) needs entirely — the slice customer may not even live nearby. Sell dinner to dinner buyers (residential geo-fencing, family bundles), not to lunch regulars. [§ 2.1 demand segmentation, § 1.1 positioning, § 10]

Q6. Delivery carries lower margins than carryout after driver costs, yet we push delivery promos. What does the promo calendar know that the P&L doesn't?

A: Nothing — the calendar runs on top-line habit and franchise-era playbooks. Run channel contribution: if delivery loses per-order, promos should push carryout (carryout-only deals) and delivery should carry its true fee. Channel-shift marketing beats channel-subsidizing marketing.

[§ 1.3 channel contribution, § 2.3, § 13]

Q7. Does owning drivers beat aggregators on net margin — re-run since insurance rates moved?

A: Re-run it: driver model = (wages + insurance + idle time + liability) vs. aggregator commission per order. Rising insurance has flipped many markets to aggregators for low-density zones while dense cores still favor owned drivers. The answer is usually hybrid — by zone and daypart. [§ 11.2

make-vs-buy, § 10 density economics, § 1.3]

Q8. Why do gourmet specialties outsell cheese on weekdays but lose on weekends — which daypart do we menu-engineer for?

A: Different buyers: weekday = individuals/pairs exploring; weekend = families with kids voting, and kids vote cheese/pepperoni. Engineer both: weekday lunch features the gourmet line; weekend bundles lead with the classics. One menu, two merchandising moments. [§ 2.1 segmentation by daypart, § 2.3, § 13]

Q9. Are coupons acquiring households or subsidizing existing ones — what does address-level redemption show?

A: The address history answers: if >60% of redemptions come from households that ordered in the prior 90 days at full price, you're subsidizing. Fence acquisition coupons to new addresses (first-order codes, new-mover mailers) and keep retention offers off public channels. [§ 8.2 fences, § 2.3, § 13]

Q10. We measure food cost monthly while cheese moves weekly. What has the lag cost over three years?

A: Compute it: theoretical vs. actual food cost variance, summed over the months where cheese spiked but prices held — typically 1–3 margin points per spike quarter. The fix is weekly theoretical food-cost reporting on your top five commodities; monthly accounting is a governance choice, not a constraint. [§ 5.1, § 1.3 variance, F1 time-indexing, § 13]

49. Ice cream & dessert shops

Q1. What is priced into the \$7 scoop: product, or the childhood memory?

A: The memory — commodity ice cream is \$3; the premium buys nostalgia, occasion, and place. That means experience consistency (line theater, generous samples, the smell of cones) IS the product, and cost-cutting the experience degrades the thing being sold. Protect experience inputs as fiercely as ingredient quality. [§ 6.1 perceived quality/aesthetics, § 8.2 blueprint, § 1.1 OrderWinner]

Q2. Our premium-ingredient story may justify the price only to us. What does churn of price-sensitive regulars show?

A: Whether the story travels: if regulars defect to cheaper shops and don't return, the ingredient narrative isn't self-funding — quality they can't taste is quality they won't fund. Test a blind comparison; if customers can't distinguish, reallocate the premium into visible experience instead.

[§ 6.1 perceived vs actual, § 2.3, § 13]

Q3. Why do birthday bookings come from neighborhoods we don't market to — what's driving them?

A: Social network propagation: one party hosts twenty kids' families from other zip codes, and the booking follows the guest list, not your ads. Track and feed the mechanism (party-guest coupons, school-network ambassadors) — your real marketing surface is the party itself. [§ 2.3 referral

mechanics, § 13 network effects, § 10]

Q4. Why do seasonal hires return yearly while year-round staff churns — what does seasonality do to loyalty?

A: Seasonal work is a ritual (summer job identity, no burnout accumulation); year-round work is a grind with retail-ceiling wages. The structural fix: build progression (shift lead, catering, winter roles) so year-round has somewhere to go, or accept churn as designed and systematize training.

[§ 4.4 job design, § 13, A9 learning curve/standard work]

Q5. When we add coffee to fight seasonality, why does it cannibalize afternoon dessert traffic instead of adding a daypart?

A: Because your coffee isn't a destination product — it attracts your existing customers earlier instead of new coffee-commuters. Counter-seasonal additions need their own demand source (morning location traffic, office density). If the location has no morning flow, winter needs a different answer (catering, wholesale, closures). [§ 2.3 counter-seasonal strategy, § 10 location, § 1.1]

Q6. Winter loses money as a ritual while catering and wholesale sit underdeveloped. What makes January a channel decision?

A: A capacity-allocation analysis: your kitchen, brand, and staff exist year-round while the storefront demand doesn't. Wholesale (grocery pints, restaurant desserts) and catering use the same assets on counter-seasonal cycles. Build the B2B channel in summer for January delivery. [§ 7 aggregate planning, § 2.3, § 4.2 utilization]

Q7. Scoops-per-hour vary 40% across staff with identical training. Does variance track tips, shifts, or the unmeasured?

A: Measure all three: usually it's shift assignment (peak staff look fast) plus individual pace. Normalize by daypart first; the residual is personal velocity — then study the fastest (positioning, pre-portioning habits) and codify. Also check if high-speed staff have lower sample generosity (a real trade-off). [§ 4.4 work measurement, § 6.3 stratification, § 13]

Q8. When grocery pints launch, brand searches rise but store traffic doesn't. Which pays the rent?

A: The store — so judge the grocery program by its wholesale margin plus its halo, not by traffic that was never its mechanism. But verify the halo exists (search lift → catering, gift cards); if the program is standalone-thin margin, it's brand advertising with a cost line. [§ 1.3 channel economics, § 2.3, § 13 attribution]

Q9. Why do most-Instagrammed flavors sell below average — which metric are we flavor-developing for?

A: For the camera: photogenic flavors (unicorn swirls, charcoal) generate earned media but often taste ordinary or alienate kids (the actual volume buyers). Keep a rotating photo-flavor as marketing budget, and develop the core menu on repeat-purchase data. Two different products with two different KPIs. [§ 6.1 aesthetics vs performance, § 13 vanity metrics, § 2.1]

Q10. Why do lines correlate with below-average per-customer spend — is the queue suppressing orders?

A: Yes: queue pressure truncates browsing and add-ons ("just a single, they're waiting"), and families with big orders avoid long lines entirely. Speed up the line (menu boards upstream, sample discipline, second register) and per-customer spend recovers — throughput and ticket size are complements here, not trade-offs. [§ 8.1 queue psychology, § 8.2, § 10 process design]

50. Juice / smoothie / boba shops

Q1. Why do cleanse participants complete the program then disappear instead of becoming daily customers?

A: Because the cleanse is a goal-completion product — finishing it closes the mental account with a sense of achievement, not a habit. The daily habit needs a different entry: post-cleanse transition programs (subscribe to 3/week maintenance) sold at the completion high, not after. [§ 13 behavioral/goal gradient, § 2.3, § 8.2 subscription]

Q2. Our busiest hours run on our least experienced staff. What if seniority ruled the 4 PM line?

A: Test it: the after-school rush is high-volume, low-complexity, speed-critical — arguably right for trained juniors; but error and waste rates at peak cost more than at lull. Optimal is mixed crews with a senior anchor at peak. Measure remake rate and ticket time by crew composition. [§ 4.4 skill matrices, § 7 shift design, § 6.3]

Q3. Why does delivery mix carry higher ratings but lower tips and worse margins — which channel do we steer toward?

A: Higher ratings because the product survives transit well and complaints never reach you; worse margins because the platform toll exceeds the tip loss. Steer margin-sensitive volume to pickup (pickup-only promos, loyalty points on direct orders); let delivery serve acquisition and convenience segments only. [§ 11.3 platform economics, § 2.3 channel steering, § 8.2 fences]

Q4. When a competitor copies our top seller cheaper, why do we keep customers three months then lose them all at once?

A: Loyalty has a grace period: habits resist one disconfirmation, but accumulated price-comparison moments cross a threshold and the switch cascades socially. The three months are your window — respond then (value bundle, loyalty lock-in, differentiation refresh), not at the cliff. [§ 13 behavioral/switching, § 2.1, § 11.4 competitive response]

Q5. Does the loyalty app increase frequency or discount existing visits — what does cohort analysis show?

A: The cohort analysis: compare member visit frequency pre/post join vs. matched non-members. If members' frequency is flat, the app is a discount distribution system for demand that existed. Redesign rewards toward frequency-building (streaks, off-peak bonuses) rather than flat percentage-back. [§ 2.2 cohort analysis, § 13, § 8.2 yield]

Q6. Why do two locations a mile apart share almost no customers — are we running two businesses with one menu?

A: Yes, and that's normal: food-and-beverage demand is hyper-local (3-minute convenience radius), so each store serves a distinct micro-market that happens to share your brand. The error is uniform marketing and menu; operate each store on its own local demand data. [§ 10 location/coverage models, § 2.1, § 1.1]

Q7. Health-positioned customers churn at 90 days; treat-positioned stay for years. Which does the menu serve?

A: The churners, if your menu leads with health claims: health goals expire (the 90-day resolve cycle), while treat habits don't. Serve both but know the economics — retain the health crowd with rotating novelty (their loyalty is to the goal), and let the treat crowd be the annuity. [§ 13 behavioral/goal cycles, § 2.1, § 2.3]

Q8. What does merchandise attach rate reveal about whether customers bought a beverage or a lifestyle accessory?

A: High attach = identity purchase (the brand as tribal signal — boba culture, wellness identity); low attach = commodity refreshment. Identity customers tolerate premium pricing and evangelize; commodity customers price-shop. The attach rate is your brand-equity thermometer. [§ 6.1 perceived quality/identity, § 13, § 1.1 positioning]

Q9. Customization slows the line 40% and drives all social mentions. Is the queue the marketing?

A: Partly, yes: the visible craft and the posted result are the brand's content engine. But the queue also caps peak revenue. Resolve by separating: keep customization theater at off-peak, and offer "fast lane" preset builds at rush. You can sell the show and the speed — to different customers, in different hours. [§ 8.1, § 8.2 yield/daypart, § 6.1]

Q10. Why does COGS swing 6 points with produce seasonality while menu prices never move?

A: Because repricing feels risky, so you eat the swing — but 6 points is the difference between profit and loss in this category. Options: seasonal menu rotation (feature what's cheap), indexed supplier contracts, or small scheduled price reviews. Static prices against volatile inputs is a policy, not a law. [§ 11.3 price fluctuation, § 5.1, § 0.3 governedBy]

51. Auto repair shops

Q1. Customers who decline recommended work return within 90 days for that exact repair. What does the estimate conversation fail to transfer?

A: Urgency with evidence: the decline means the customer didn't believe the consequence timeline. Digital inspection photos plus a plain-language consequence ("this fails completely within ~3 months and strands you") transfer what a line item can't. Track decline-to-return rate as your estimate-quality KPI. [§ 6.3 SERVQUAL assurance, § 8.2 blueprint, § 13]

Q2. Does the loaner fleet buy retention or unmeasured goodwill — what does retention of users vs non-users show?

A: Run the comparison: if loaner users retain 15+ points higher, the fleet is a retention asset — cost it against LTV. If retention is equal, it's goodwill decoration. Also check the alternative: shuttle service or ride credits may buy the same retention at half the capital. [§ 1.3 BalancedScorecard customer perspective, § 4.1 capital, § 2.3]

Q3. Why do comebacks cluster by parts brand rather than technician — and who is watching?

A: Nobody, apparently — this is exactly what a parts-quality log catches. Brand-clustered comebacks are special-cause variation pointing at supplier quality (white-box parts). Track comeback rate by parts brand × job type, then renegotiate or respect your sourcing. [§ 6.3 SPC special cause, § 11.2 supplier management, A8]

Q4. What does retention look like two quarters after a service advisor exits — which relationships survive?

A: Typically a visible dip: the advisor was the trust interface, and customers followed their calls. Survivors are customers with multi-touch relationships (techs, owner). Build depth deliberately: introduce customers to the team, keep communication in shared systems, so the relationship is with the shop. [§ 13 relationship capital, § 11.4 key-person risk, § 8.2]

Q5. Best reviews describe communication; worst describe surprise. Which does the estimate process create?

A: Both — the estimate process is the review factory. Communication = proactive updates and approved-work clarity; surprise = unapproved additions and price drift at pickup. Standardize the update cadence and the no-work-without-approval rule; reviews will follow the process, not the luck. [§ 8.2 blueprint, § 6.3 SERVQUAL responsiveness, § 13 standard work]

Q6. Are maintenance reminders bringing customers back to us or just reminding them the car needs service?

A: Check capture rate: reminder-to-booking attribution. Generic reminders ("time for service") educate the market and lose the customer to whoever's convenient; effective ones are specific (vehicle, service due, one-click booking with your shop). Measure capture by reminder design.

[§ 2.3, § 12 CRM, § 13]

Q7. Why do we diagnose intermittent electrical faults at a loss as a matter of pride — is there a profitable pricing model?

A: Yes: diagnostic-time billing, stated upfront ("electrical diagnosis billed in 30-min increments, applied to repair"). The pride is real capability — the loss is pricing it as a flat gamble. Customers accept diagnostic pricing when framed as expertise; only shops assume they won't. [§ 2.3 pricing structure, § 1.1 OrderWinner expertise, § 13]

Q8. Digital inspections raise approval rate but lower average ticket. Which do we want?

A: Both metrics are proxies — what you want is lifetime trust. Higher approval + lower ticket usually means customers approve safety items and defer the rest (good: honest triage) instead of feeling bundled into big tickets (bad: breeds distrust). Optimize approval rate and 12-month retention, not ticket size. [§ 13 Goodhart, § 6.3 assurance, § 1.3 BalancedScorecard customer perspective]

Q9. Why do fleet accounts negotiate our rate down, then generate our steadiest profitable hours?

A: Because utilization is the hidden variable: fleet work fills bays at off-peak with no marketing cost, no waiting-area friction, and predictable payment — the rate is lower but the effective margin per bay-hour is high. This is legitimate yield management; just floor the rate at true cost. [§ 8.2 yield, § 4.2 utilization, § 1.3]

Q10. Our effective labor rate runs 20% below posted. Which write-downs does nobody own?

A: The invisible ones: goodwill discounts, absorbed diag time, "while you're in there" extras, and tech-time overruns on flat-rate jobs. Effective rate = posted $\times$ (billed hours $\div$ actual hours) $\times$ (collected $\div$ billed). Assign each term an owner; most shops find the 20% splits across all three.

[§ 1.3 rate realization, § 0.3 measuredBy, § 13]

52. Car washes

Q1. Why do detail upsells convert at the pay station but not from highest-tenure members?

A: Members have routinized their purchase (the unlimited plan is the settled decision); the pay-station upsell catches non-members mid-decision. Members need a different mechanism: periodic member-exclusive detail offers, not point-of-sale pitches. Same customer, different buying state. [§ 13 behavioral/habit, § 2.3, § 8.2 fences]

Q2. When we raise membership prices, new signups hold while grandfathered members complain and stay. Which group are we pricing for?

A: You're discovering that existing members have high switching inertia (complaint $\neq$ churn) and new signups are price-anchored to current value. Price for new acquisition normally; for the base, use small annual increases with notice — the data says they stay. [§ 2.3 pricing, § 13 behavioral/inertia, § 8.2]

Q3. Why does membership churn spike in month two, before the value habit could form?

A: Buyer's remorse plus first-bill shock: month two is when the second charge posts against a member who used the wash twice. Intervention window is days 7–30: usage nudges ("you have unlimited washes — here's a free tire shine"), not cancellation saves. [§ 13 behavioral, § 2.1 churn timing, § 8.2 subscription design]

Q4. Does the free-vacuum policy drive volume or just lengthen the queue — what does per-hour throughput show?

A: Check whether vacuums are the binding constraint at peak: if the tunnel waits on vacuum bays backing up, free vacuums are taxing your conveyor capacity. Options: vacuum-time limits, paid express lanes bypassing vacuums, or staff-assisted vacuum at peak. [A3 bottleneck, § 8.1, § 10 layout]

Q5. When it rains, membership revenue holds but retail vanishes. Which one built the business plan?

A: If the plan assumed retail (per-wash) volume, rain exposes the flaw; if it assumed membership, rain is the sales pitch ("wash whenever, weather-proof"). The industry's trajectory says membership is the plan — retail becomes the acquisition funnel for it. Align the P&L model to the membership reality. [§ 2.1 demand weather-sensitivity, § 8.2 subscription, § 1.1]

Q6. Why do chemical costs per car rise with volume — opposite of the supplier's pitch?

A: Because dilution and dosing drift: high volume stresses calibration (pumps wear, settings get "adjusted" by attendants chasing cleaning quality), and reclaim systems underperform at throughput peaks. Audit dosing calibration weekly at volume — the supplier's pitch assumed lab conditions. [§ 6.3 process control, § 5.1 consumables, § 4.3 maintenance]

Q7. Our site selection bet: convenience business or subscription business. Which did the traffic data say we made?

A: Read your own mix: high capture of pass-by commuters = convenience site (retail pricing matters); high membership penetration = subscription site (density of rooftops matters). The traffic data decides which P&L physics you're running — mismatch (convenience site pushing memberships to transients) is the failure mode. [§ 10 location models, § 2.1, § 1.1]

Q8. Why does the newest location cannibalize the oldest at the transfer rate our model said was impossible?

A: Because the model assumed trade areas don't overlap, and members are mobile: unlimited members rationally migrate to the newer/closer site. Cannibalization was always computable from member home/work zip overlap — the model just never got that data. Before site three, map member geography. [§ 10 location/cannibalization, § 2.2 data, § 13 assumption auditing]

Q9. Unlimited members wash less over time but never cancel. When is that our margin, and when our lawsuit?

A: It's margin while members feel they could use it (gym-membership economics); it becomes legal/reputational risk when cancellation is obstructed or value claims turn deceptive. Keep cancellation one-click easy and occasionally prompt usage — dormancy is your margin, but entrapment is your liability. [§ 13 behavioral ops, § 14 compliance, § 8.2]

Q10. Why do attendants' upsell rates vary 5× with the same script — script or lane assignment?

A: Separate the factors: swap attendants across lanes for a week. If rates follow the person, it's delivery (confidence, timing, reading the customer); if they follow the lane, it's customer mix. Most likely person — then codify the top performer's actual behavior (usually: offer one item, name the benefit, shut up). [§ 6.3 stratification, § 13 standard work, § 4.4]

53. Used car dealerships

Q1. Are we in the car business or the financing business — which survives a cash-buyer market?

A: Check gross per unit: if finance reserve + products exceed front-end gross, you're a finance business with a car lot. In a cash-buyer market that model starves. Neither is wrong — but inventory selection, pricing, and sales process must match the real business. [§ 1.1 emergent strategy, § 1.3, § 13]

Q2. Why do online leads close at half the walk-in rate at a quarter of the cost — which funnel does the floor work?

A: The floor works walk-ins (visible, immediate), letting online leads age in the CRM — which then "proves" they don't close. Self-fulfilling: online leads need 15-minute response and a different process (appointment-setting, not car-selling). Staff the internet funnel like it matters; the cost math says it does. [§ 13 self-fulfilling metrics, § 8.2 blueprint, § 1.1 OrderWinner speed]

Q3. Why do repeat buyers return every 3–4 years with zero contact in between?

A: Because no one owns the ownership cycle: the sale ends the relationship operationally, and 36 months later the customer is back in the open market. A lifecycle contact program (service reminders, equity position updates, trade-in offers at loan maturity) captures the return you currently get by luck. [§ 2.3 lifecycle demand, § 12 CRM, § 13]

Q4. Our best salesperson has lowest satisfaction scores and highest repeat-referral rate. Which predicts return buyers?

A: Repeat-referral: satisfaction surveys measure the sales-day experience; repeat behavior measures the ownership outcome. The salesperson probably pressures at close (bad surveys) but sells the right car at a fair price (buyers return). Decide which you optimize — but know the survey is measuring the shallower thing. [§ 13 Goodhart, § 6.3 vs § 1.3 BalancedScorecard customer perspective, § 6.1]

Q5. Why do 30-day warranty claims cluster on vehicles we certified hardest — what does certification inspect?

A: The checklist, not the car: certification is a documentation process vulnerable to pencil-whipping under lot pressure, and "certified" cars are often the ones pushed hardest (price premium, faster turn). Claims clustering means inspection is theater at your worst moments. Audit: re-inspect a sample of certified units. [§ 6.3 acceptance sampling, § 13 Goodhart, § 6.2 external failure]

Q6. When rates move, why does our inventory mix hold for six months — what would rate-responsive buying look like?

A: Because acquisition runs on habit and auction relationships, not demand signals. Rate-responsive: when rates rise, payment-sensitive buyers shift down-market — buy cheaper, faster-turning units; when rates fall, stretch into higher-GPU inventory. Tie the buy plan to rate-triggered affordability bands. [§ 2.2 leading indicators, § 5.1 inventory mix, § 11.2]

Q7. Does the service department exist to profit or feed the lot — which does its pricing say?

A: If internal reconditioning gets priority and retail service is priced indifferently, it's a cost center feeding the lot. That's legitimate — but then don't judge it on retail-service KPIs, and don't neglect retail customers whose service experience seeds their next purchase. Name its role; measure it accordingly. [§ 1.1 structural role clarity, § 13 Goodhart, § 1.3]

Q8. Why do highest-gross units sit longest — does pricing cause the aging or vice versa?

A: Pricing causes it: high gross asks extend days-on-lot, aging then forces the discount that erases the gross. The market-clearing price at day 10 beats the aspirational price at day 60 once floorplan costs count. Set aging-based price ladders (auto-reductions at day 21/35/45). [§ 5.1 inventory aging, § 8.2 yield/dynamic pricing, § 4.1 carrying cost]

Q9. When customers negotiate, discounts come from front-end gross while finance income stays untouched. Who taught them?

A: The internet and the industry itself: decades of invoice-price disclosure made front-end transparent while finance stays opaque — so negotiation concentrates where information lives. Your defense is product bundling and finance-menu transparency that reframes total cost, not line-item siege. [§ 13 information asymmetry, § 2.3 price architecture, § 6.1]

Q10. Auction purchases beat trade-ins on margin but lose on reconditioning surprises. Which wins all-in?

A: Total it per channel: purchase cost + recon (actual, including surprises) + floorplan days-to-frontline. Trade-ins typically win all-in (you see the car's history, recon is predictable) while auctions win on selection. The all-in number usually rebalances the buy mix toward street purchases. [§ 1.3 total-cost analysis, § 11.2 sourcing, § 2.1]

54. Auto detailing

Q1. Why do dealership contracts give steadiest volume and worst effective hourly rate — what would exiting cost?

A: Dealer work is margin-thin by design (they capture the channel value) but it fills capacity between retail jobs. Exit cost = lost utilization × contribution, minus freed capacity for retail. Usually the right move is capping dealer share, not exiting — it's your base-load, not your business.

[§ 8.2 yield/base-load logic, § 4.2 utilization, § 1.3]

Q2. Does the photo library sell the work or set unreachable expectations — which complaints reference it?

A: Check complaint content: if "didn't look like the photos" appears, your library shows best-case outcomes (perfect paint, studio light) — acquisition assets that become delivery liabilities. Show representative results and label conditions; expectation management is cheaper than disappointment recovery. [§ 6.3 communication gap, § 6.1 perceived quality, § 13]

Q3. Which framing survives the customer's spouse: time saved, transformation, or protection?

A: Protection, typically: "preserves resale value and protects the paint" justifies the spend to the household CFO; "it'll look amazing" invites "it looks fine." Lead the quote with protection economics; let transformation be the emotional closer. [§ 13 framing/household decision, § 2.3, § 6.1 durability]

Q4. Why do two top detailers with identical quality differ 30% in daily revenue — what does the faster one skip?

A: Observe: usually upsell conversation time (the higher earner sells add-ons mid-job) or rework avoidance (checklists, no comebacks). If genuinely faster at identical quality, they've got better motion economy — film it and teach it. Either way the delta is transferable knowledge. [§ 4.4 work measurement, § 13 standard work, § 2.3]

Q5. When customers supply their own products, why do we take the job — and what does the rework rate show?

A: The rework rate justifies a policy: customer products (wrong chemistry, unknown dilution) fail at 2–3× your standard jobs, and the failure lands on your reputation. Either decline, charge a client-product surcharge with warranty carve-outs, or convert them ("here's what we use and why"). [§ 6.2 external failure risk, § 0.3 governedBy, § 11.2 input control]

Q6. Ceramic customers refer at 4× wash customers, yet marketing spend goes to wash packages. Which customer is the budget for?

A: The ceramic customer: high ticket, high advocacy, and the referral carries a price-insensitive introduction. Wash packages are the funnel, not the destination. Rebalance: use wash marketing to acquire, then engineer the wash-to-ceramic upgrade path (paint assessment at delivery). [§ 2.3 product ladder, § 1.3 BalancedScorecard customer perspective, § 13]

Q7. Why do spring bookings fill eight weeks out while winter capacity idles — what service haven't we invented for January?

A: The winter-specific one: salt/grime decontamination, interior deep-clean after holiday hauling, undercarriage protection. Winter demand exists but your menu is summer-framed. Package "winter recovery" and sell it in November as a pre-booked slot. [§ 2.3 counter-seasonal demand, § 7 capacity smoothing, § 13]

Q8. Why does mobile margin beat shop margin while shop occupancy stays fixed — which is the real business?

A: Mobile, by the numbers — but the shop enables premium services (coatings need controlled environments). The real structure: mobile is the volume engine, shop is the premium lab. Run the shop's fixed costs against premium-only work and push everything else to mobile. [§ 1.2 product-process fit, § 10 facility, § 1.3]

Q9. When we offer maintenance plans after full details, why do highest-spend customers decline most?

A: Because they just spent \$400 and the plan reads as a subscription pitch at the moment of maximum spend — timing, not value. Offer at the second touchpoint (the follow-up call when they rave about the car), or embed a first maintenance visit in the premium package. [§ 13 timing/framing, § 2.3, § 8.2 blueprint]

Q10. Why do five-star reviews mention "on time, returned calls" rather than the paint correction we pride ourselves on?

A: Because customers can't evaluate paint correction — they evaluate the operational courtesies they understand. Basic reliability is your review currency; craft is invisible except to enthusiasts. Systematize the basics (confirmations, on-time starts, proactive updates) — they're process outputs, cheaper than craft. [§ 6.3 SERVQUAL reliability/responsiveness, § 6.1, § 13]

55. Tire shops

Q1. Why do fleet tire accounts churn on price while our records show downtime savings they never count?

A: Because the savings live in their operations data, not their procurement spreadsheet — procurement compares price per tire, period. Sell the accounting: quarterly business reviews translating your service records into their downtime dollars make the premium defensible. Unquantified value gets shopped. [§ 6.1 perceived quality, § 11.2 supplier value documentation, § 13]

Q2. Tire margin compresses while service attach stays flat, and the model depends on attach. When does the tire sale stop being worth winning?

A: When tire gross margin no longer covers the bay-time opportunity cost plus acquisition discount — i.e., when the tire sale stops generating the service attach it exists to feed. Compute margin per bay-hour with and without the tire line; if attach can't lift it, the tire becomes a traffic product you should reprice or de-emphasize. [§ 1.3, A2 opportunity cost, § 2.3 loss-leader logic]

Q3. When alignment machines sit idle on Saturdays — our busiest day — what scheduling assumption made that hole?

A: The assumption that alignment customers book weekday appointments: your Saturday walk-in tire buyers need same-day alignment, but the tech certified for the machine works Monday–Friday. Cross-train weekend staff or schedule the alignment tech on Saturday — the idle machine on your busiest day is the demand signal you're ignoring. [§ 4.4 skill matrices, § 7 scheduling, § 4.2 utilization]

Q4. Why do customers decline nitrogen and TPMS services that take minutes — price or explanation?

A: Explanation: "nitrogen fill, \$20" sounds like air-for-money; "stable pressure means even wear and fewer TPMS lights" sounds like tire-life insurance. Low-attach micro-services die from jargon pricing. Script the benefit in tire-longevity dollars and watch attach move. [§ 6.3 assurance, § 13 framing, § 2.3]

Q5. Why do road-hazard claims run above actuarial rate at one location — what is that service writer doing?

A: Writing claims generously to delight customers (or worse) — claim adjudication discretion varies by writer, and one location's culture normalized loose approval. Above-actuarial claims are a controlled experiment revealing your leakage: standardize claim criteria and audit that store's approvals. [§ 6.3 special cause, § 0.3 governedBy, § 13]

Q6. Customers price-shop tires by phone but buy alignment and brakes on trust. Which service is the phone script written for?

A: Probably the tire shopper (the visible demand), but the script's job is conversion to visit, not tire price victory — once the car is on the lift, trust services carry the margin. Rewrite the script: competitive tire quote fast, then pivot to the inspection that surfaces the trust work. [§ 1.1 OrderWinner by product, § 2.3, § 8.2 blueprint]

Q7. Why does inventory carry sizes that sold three years ago while the fastest-growing fitment is on backorder?

A: Because buys follow history without trend adjustment: fitment mix shifts as the local vehicle parc turns over (trucks/EVs upsizing wheels). Run ABC with trend: A-items by 12-month velocity with a trend factor, and cull the zombie sizes eating your tire racks. [§ 5.2 ABC classification, § 2.2 trend forecasting, § 5.1]

Q8. Does gross profit come from rubber or wrenches — why does hiring budget assume the former?

A: Run the split: typically service labor carries 60%+ of gross profit while hiring chases tire-busting generalists. If wrenches pay, hire and retain technicians accordingly (and price tires to feed the service bays). The budget confesses a tire-business identity in a service-business P&L. [§ 1.3, § 1.1 emergent strategy, § 4.4]

Q9. Does free rotation bring customers back to us or teach them to buy tires anywhere?

A: Check conversion: what fraction of free-rotation visits bought their tires from you vs. elsewhere? If elsewhere dominates, you're running a free service for competitors' customers. Fence it: free rotations with purchase, paid otherwise — the rotation is your retention device, not a public good. [§ 8.2 fences, § 2.3, § 13]

Q10. When we match online tire prices, why do we win the sale but lose install margin to a "booking fee" nobody questions?

A: Because the install fee is where the transaction's remaining margin hides, and discounting it to close the match-sale concedes the only profit left. Set a floor: matching online prices is fine only with install at full rate — otherwise you're paying customers to use your equipment. [§ 2.3 price architecture, § 1.3, § 13 anchoring]

56. Collision / body shops

Q1. When customers choose us over the insurer's recommended shop, what did they hear — can we bottle it?

A: Interview them: usually "you work for me, not the insurance company" — independence, OEM parts advocacy, and someone fighting for full repair. Yes, bottle it: that's your consumer marketing message and your estimate-script identity. The DRP-alternative positioning is a deliberate product. [§ 1.1 MarketPositioning, § 6.3 assurance, § 13]

Q2. What does the insurer scorecard actually measure — quality or paperwork fluency — and which do we train for?

A: Paperwork fluency: cycle time, estimate compliance, photo documentation, supplement discipline — the administrable metrics. Quality rides along unmeasured. Train for both but know the scorecard's nature: it optimizes their cost process, not the repair. Your reputation with customers is the counterweight asset. [§ 13 Goodhart, § 11.3 coordination, § 6.3]

Q3. When a body tech becomes a painter, why do quality metrics transfer but speed resets to zero?

A: Because quality is judgment (transferable craft sense) while speed is motor-skill and process familiarity (not transferable). A new painter is on a fresh learning curve — plan it: reduced productivity quota during transition, mentor pairing, and don't measure them against journeyman speed for 6 months. [A9 learning curve, § 4.4 skill matrices, § 13]

Q4. Why does cycle time improve in months when supplements worsen — is speed bought with thoroughness?

A: Yes: rushing the blueprint misses hidden damage, so supplements spike after teardown — trading total cycle for measured phase speed. The metric sees keys-to-keys; the truth is in supplement timing and count. Measure first-supplement rate alongside cycle time to force honest blueprints. [§ 13 Goodhart, § 6.3, § 9 process design]

Q5. Does OEM certification pay in rate or volume — which has it delivered?

A: Audit referral flow: certifications pay in volume only if insurers/OEM programs actually steer work to you (check certified-shop referral data); the rate premium rarely materializes. If neither materialized, certification is a qualifier (table stakes for certain work) — renew accordingly, not aspirationally. [§ 1.1 OrderQualifier, § 1.3 ROI, § 11.2 channel]

Q6. DRP relationships fill bays and compress margins. What would the book look like minus the worst one?

A: Model it: worst-DRP revenue $\times$ its margin vs. replacement volume from customer-pay/retail channels at higher margin, including the marketing cost to get it. Most shops find dropping the worst DRP frees capacity that better work refills within a quarter. Run it before renewal. [§ 11.2 channel economics, § 1.3, § 8.2 yield]

Q7. We track touch time but not the two-day gaps between touches, where cycle time lives. What would a gap log show?

A: The real bottleneck map: parts waits, approval waits, sublet scheduling, and tech reassignment — the car sits in [Starved/Blocked](#) states invisible to touch-time metrics. One month of gap logging typically shows 60–80% of cycle time is waiting, and the fixes are scheduling and parts process, not faster techs. [§ 3.4 process states, A1 Little's Law, § 7]

Q8. Why do total-loss decisions cluster on estimates from one adjuster relationship — what does that do to revenue mix?

A: That adjuster totals borderline cars (their severity philosophy), converting your potential repairs into nothing. Revenue mix shifts silently through one relationship. Diversify adjuster exposure across carriers and track total-loss rate by adjuster — one person's philosophy shouldn't govern your bay loading. [§ 11.4 concentration risk, § 2.1, § 13]

Q9. Why do parts margins vary by estimator on identical insurer rate schedules?

A: Because margin lives in sourcing choices: OEM vs. aftermarket vs. LKQ selection, vendor discounts, and return discipline — all estimator decisions within the same rate schedule. Standardize the sourcing hierarchy (with documentation for exceptions) and the variance collapses. [§ 11.2 sourcing structure, § 13 standard work, § 1.3]

Q10. Why do customer-pay jobs produce best reviews and margins while intake is built for insurance work?

A: Because customer-pay buyers chose you (trust, advocacy) and pay full rates without insurer suppression — while your intake treats them as walk-in anomalies. Build the retail lane: retail-oriented estimating, payment options, and marketing to post-warranty vehicle owners. It's your highest-quality demand hiding in a DRP-shaped process. [§ 1.1 focus/positioning, § 8.2 blueprint, § 2.3]

57. Oil change shops

Q1. Why does the 10-minute promise hold at 9 AM and fail at noon — which step eats the time?

A: The queue, not the bay: noon arrivals bunch (lunch-hour demand) while bay count is fixed — the 10 minutes becomes 10 + wait. The promise fails at arrival rate, not service rate. Smooth arrivals (mid-day pricing, appointment slots) or staff the lunch surge. [§ 8.1 M/M/c queues, § 2.1 arrival patterns, § 7]

Q2. Does the fluid-exchange menu exist for vehicle health or ticket size — what do our own techs' cars get?

A: Ask the techs: if their own cars skip the exchanges, the menu is ticket inflation and every sale erodes trust capital. Keep services with OEM-schedule justification, drop or reframe the rest. In a repeat business, the menu's credibility is the asset. [§ 6.1 perceived quality/integrity, § 13, § 2.3]

Q3. Do reviews reward speed or not-overselling — which does the pay plan pay for?

A: Check the texts: reviews cite "in and out fast" and "didn't push stuff" — while pay plans typically reward ticket size (upsells). If pay rewards what reviews punish, you've structured a slow trust-erosion. Pay on cars-per-hour + audit-verified necessary upsells. [§ 13 Goodhart, § 6.3 SERVQUAL, § 1.3]

Q4. Why does manager turnover predict customer defection two quarters later, where customers never learn the manager's name?

A: Because the manager is the operating system: their departure degrades speed, quality consistency, and upsell discipline — customers defect from the experience drift, not the person. The two-quarter lag is the decay curve. Treat manager transitions as service-continuity events with overlap and audits. [§ 13 institutional knowledge, § 6.3 process consistency, § 2.2 leading indicator]

Q5. Why do coupon customers return without coupons at 20% while full-price first-timers return at 45%?

A: Selection: couponers are deal-loyal, not shop-loyal — they follow the next coupon. Full-price first-timers chose you on proximity/reputation and stay for it. Coupons are fine for filling idle capacity, but measure retained margin, not redemption; the coupon cohort often never becomes profitable. [§ 2.1 segment selection, § 8.2 fences, § 13]

Q6. When bays max out, why do we add marketing instead of a bay — what's the constraint math?

A: Backwards: marketing feeds demand into a saturated bottleneck — throughput can't rise, so spend converts to wait times and defections. The math: if bays run >85% utilization at peak, the next dollar belongs in capacity (bay, techs, hours) or demand shaping, not acquisition. [A2/A3, A4, § 1.1 StructuralDecisions]

Q7. Customers arrive with internet research on synthetic vs. conventional; our recommendation wins half the time. Which half, and why?

A: The half where your recommendation matched the research: when you agree, you win; when you push against (upsell synthetic on an old beater, or conventional against their research), you lose the sale and some trust. Align recommendations with OEM specs and explain the why — credibility converts better than conviction. [§ 6.3 assurance, § 13, § 2.3]

Q8. Why do we lose customers to dealer quick-lanes exactly when warranties expire — what's our never-made counter-offer?

A: Because the warranty period trained them to the dealer, and expiration is an unmarked transition. The counter: acquire them before expiry — "warranty-expiration inspection" offers, records transfer, price comparison versus dealer service menus. The switch moment is predictable; your silence makes the dealer the default. [§ 2.3 lifecycle triggers, § 1.1, § 13]

Q9. Why do upsell attach rates vary fourfold by technician — persuasion or honesty?

A: Audit against vehicle condition data: if high-attach techs sell services the vehicle records support, it's persuasion skill; if their attach items show no pattern of need, it's aggressive selling that will surface as distrust churn. Standardize inspection-driven recommendations and the variance becomes visible as either skill or noise. [§ 6.3 stratification, § 13, § 6.1 integrity]

Q10. Why do sticker reminders underperform text reminders, yet we keep both unmeasured?

A: Habit — stickers are legacy infrastructure nobody kills. Measure cost-per-returned-customer per channel: texts win on timing precision and click-to-book; stickers win nothing measurable. Kill or keep based on the number, and reinvest in the channel with attribution. [§ 0.3 measuredBy, § 12 CRM, § 13 Goodhart]

58. Towing services

Q1. Average response looks great; the 90th percentile is losing us the rotation slot. Which number does the contract measure?

A: The tail — rotation contracts measure worst-case reliability because their promise to the public is dependability, not averages. Manage the tail: surge coverage protocols, mutual-aid agreements, and staging during known peaks. Averages are for marketing; P90 is for contracts. [§ 6.3 distribution thinking, § 1.1 OrderWinner dependability, § 8.1]

Q2. Our best drivers leave to start one-truck operations. What would keeping them cost vs. replacing them?

A: Compare: raise + retention bonus vs. recruitment + training + the contracts their departure rattles. But the structural answer: they leave for autonomy and upside, not just money — offer a profit-share truck or senior-driver equity track before you lose another operator to entrepreneurship. [§ 13 incentives, § 4.4 job design, § 1.3 turnover cost]

Q3. When fuel spikes, contract rates hold for a year while cash rates adjust weekly. Which book grows those years?

A: The cash book, rationally — but your capacity allocation then drifts toward cash calls, degrading contract response times and risking the rotation slot that's your base-load. Add fuel escalators to contracts at renewal; without them, every spike year taxes your contract relationships. [§ 11.2 contract design, § 11.3 price fluctuation, § 8.2 capacity allocation]

Q4. Does the roadside-membership math work if members actually use it — what does utilization by tenure show?

A: Check the curve: memberships profit on low utilization (breakage); if long-tenure members learn to use it (battery jumps, lockouts), the product inverts into a loss. Cap uses per year or price tiers by usage — a membership that punishes engagement is a lawsuit-shaped product. [§ 8.2 subscription economics, § 13 breakage, § 14 compliance]

Q5. Police-rotation calls carry worst margins and best volume. What happens if we exit the rotation?

A: You lose the base-load that keeps trucks and drivers productively busy between cash calls — and the credibility that feeds municipal and cash work. Model the exit honestly: freed capacity × realistic refill rate. Usually the answer is renegotiate terms or re-bid scope, not exit. [§ 8.2 base-load/yield, § 4.2 utilization, § 11.2]

Q6. Why do impound/storage fees — our highest-margin line — generate our only regulatory risk?

A: Because the margin comes from involuntary customers with statutory protections: fee caps, notice requirements, disposition rules. High margin + hostile customer + regulation = scrutiny. Treat impound as a compliance product first (process perfection, documentation) — the margin survives audits only if the process does. [§ 14 compliance, § 11.4 regulatory risk, § 0.3 governedBy]

Q7. Follow the software budget: dispatch company that owns trucks, or trucking company that answers phones?

A: The budget confesses the truth: towing economics are won in dispatch (call capture, ETA accuracy, driver utilization) — the truck is a commodity asset. If software spend lags truck spend by 10×, you've misallocated toward the depreciating asset and away from the differentiating system. [§ 12 platform layer, § 1.1 StructuralDecisions, § 13]

Q8. Why do the last three calls of a 14-hour day produce the most damage claims?

A: Fatigue: degraded judgment in loading, clearances, and securing — damage claims are your fatigue meter. Hours-of-service discipline isn't just regulatory; it's claims prevention. Cap shift length or rotate drivers on long days; the claims data already priced fatigue for you. [§ 14 safety/fatigue, § 4.4 ergonomics, § 6.2 external failure]

Q9. Why do repair-shop referrals send their worst-margin calls — which shop sends the good ones?

A: Because shops offload the problem calls (no-start diagnostics, angry customers, distant drops) and keep the profitable tows for their preferred partner. Analyze margin by referring shop: reward the shops sending good work (faster response, co-marketing), renegotiate or deprioritize the dumpers. [§ 11.2 channel management, § 1.3, § 13]

Q10. Why do motor-club calls generate double the complaints of cash calls at identical service?

A: Expectation asymmetry: club members were promised "free" 30-minute rescue by the club's marketing; cash customers negotiated reality with you directly. The club sets expectations you can't meet at club rates. Push complaint data back to the club at contract time — their promise, your reputation. [§ 6.3 SERVQUAL expectation gap, § 11.3, § 2.3]

59. Transmission shops

Q1. Why are reviews bimodal — five stars or one star, nothing between?

A: Because the job is binary: a \$3,500 invisible repair either works (gratitude) or fails/under-delivers (betrayal). One-stars come from expectation gaps — cost shock, recurrence, time. The middling experience doesn't exist in transmission work. Manage the two moments that create one-stars: the authorization conversation and the warranty callback. [§ 6.1 perceived quality, § 6.3, § 8.2 recovery]

Q2. Why do general shops send us failures but keep the services — what would change the flow?

A: Because services (fluid exchanges) are profitable and safe; failures are liability they'd rather hand off. Change the flow by rewarding the whole referral: pay referral value on services too, or bundle ("send us the exchange work at wholesale and we'll keep failures priority"). [§ 11.2 channel incentives, § 11.3, § 2.3]

Q3. Does the free tow acquire customers or breakdowns nobody wanted — what does towed-vs-driven profitability show?

A: Run it: towed-in jobs skew to catastrophic failures (high ticket but high discount pressure and payment risk) while drive-ins include the profitable service work. If towed jobs' collection rate and margin lag, the free tow is buying adverse selection. Fence it: free tow with approved repair over a threshold. [§ 2.1 adverse selection, § 8.2 fences, § 1.3]

Q4. Are we in the transmission business or the trust business — which does our warranty honor in practice?

A: The trust business: no customer can verify a rebuild — they buy your warranty and your manner. So the warranty's real-world behavior (honored cheerfully vs. lawyered) IS the product. Shops that honor warranties loosely win the referral game that is this entire market. [§ 6.1 perceived quality/assurance, § 6.3, § 1.1 OrderWinner]

Q5. Why does the lead builder's output set the shop ceiling — and why hasn't their diagnostic process become a checklist?

A: Because their diagnostics are pattern-recognition built over years — tacit knowledge that resists checklists. But it partially codifies: video their teardown narration, build decision trees for the common failure modes, and pair apprentices. The ceiling stays until the knowledge becomes process. [§ 13 tacit knowledge/communities of practice, A9, § 4.4]

Q6. Customers authorize a \$3,500 rebuild but balk at a \$250 diagnostic. What do they believe they're buying?

A: A thing (the rebuild) versus air (the diagnosis): customers pay for tangible transformation, not information — even when the information determines the right transformation. Reframe: diagnosis as "confirmed written quote with teardown findings" bundled into the repair. Sell certainty, not inspection. [§ 13 tangibility bias, § 2.3 framing, § 6.3]

Q7. When the car is worth less than the repair, why do half proceed — are we advising or selling?

A: They proceed for reasons the spreadsheet misses: no car payment, known history, no better option at that cash level. Your job is honest framing (total cost of alternatives) — customers who proceed with full information become evangelists; those who feel sold become one-star reviews. Advise; the trust pays longer. [§ 13, § 6.3 assurance, § 6.1]

Q8. Why does quote-to-approval drop when we present options (used/rebuilt/reman) instead of one recommendation?

A: Choice overload: three options transfer a technical decision to a customer unqualified to make it — they freeze, "think about it," and shop it around. Present one recommendation with the rationale; keep alternatives in reserve for objections. Confidence sells; menus stall. [§ 13 choice architecture, § 6.3, § 2.3]

Q9. Why do our two reman suppliers have 4× different failure rates at identical prices — and why haven't we switched?

A: Because failure data lives in your warranty claims, not the purchase order — and nobody connected them. The 4× difference costs you comebacks, labor, and reputation, dwarfing any price delta. Track failure rate by supplier SKU and switch; inertia here is expensive. [§ 11.2 supplier scorecards, § 6.2 external failure, A8]

Q10. Rebuild warranty claims cluster at months 11–13. Workmanship, parts, or driving behavior we never discuss?

A: Stratify by failure mode: early-life (0–3 mo) = workmanship; mid = parts; the 11–13 cluster suggests break-in abuse and fluid-neglect — the customer behaviors you never contracted for. Add a 500-mile checkup requirement and break-in instructions to the warranty terms; claims will drop. [§ 4.3 bathtub curve, § 6.2, § 0.3 specifiedBy warranty terms]

60. Mobile mechanics

Q1. Does our no-shop-overhead advantage survive honest drive-time accounting — what's our real hourly rate by zip?

A: Compute it: billable hours ÷ (work + drive + parts runs) per zip cluster. Mobile models win on dense routes and lose on sprawl — your real rate by zip will show a map of where you're a business and where you're a charity. Price or zone accordingly. [§ 10 routing/density, § 1.3 true costing, § 4.2 utilization]

Q2. Why do we decline the jobs (transmissions, alignments) our best customers eventually need, sending them to shops that keep them?

A: Because your capability boundary becomes their switching point: the customer who leaves for a transmission job experiences a shop that does everything, and never returns. Counter: formal referral partnerships with revenue share or reciprocal arrangements — lose the job, keep the customer. [§ 11.2 network strategy, § 2.3 retention, § 1.1 scope decisions]

Q3. Customers approve work by text at higher rates than in person. What pressure does the driveway remove?

A: Social pressure: in person, declining feels like confrontation with the person standing there — so customers stall ("let me think"); by text, the decision is private and quick. Text also creates documentation. Make text-with-photos the standard approval channel; it converts better and records everything. [§ 13 behavioral, § 8.2 blueprint, § 12]

Q4. When weather cancels a day, revenue vanishes but goodwill holds. Which does our pricing assume?

A: Pricing assumes full availability — you're absorbing weather variance as uncompensated downtime. Options: weather-day rescheduling priority (retention), a modest margin premium as weather insurance, or indoor-work partnerships for bad days. The variance is real; price it or buffer it. [A4 buffer choice, § 11.4 weather risk, § 2.3]

Q5. We cap at four jobs daily while shop techs do eight. How much is physics vs. scheduling discipline?

A: Measure the day: drive time, parts runs, and setup/tear-down per site are physics (2–3 jobs' worth); the rest — gaps, unplanned supply runs, poor sequencing — is discipline. Tight routing and van-stock discipline can buy job five; job eight requires a shop. [§ 4.2 capacity decomposition, § 10 routing, § 5.1 van stock]

Q6. Why does scaling to two techs halve our own income — which growth assumption was wrong?

A: The assumption that revenue scales with techs while you keep your billable role: adding a tech converts you from producer to manager/dispatcher (unbillable), and their margin must cover their wage plus your lost production. The fix: either stay boutique-solo or hire past the valley (3+ techs with you purely managing). [§ 1.1 StructuralDecisions, A3, § 1.3]

Q7. Why do fleet accounts adopt fastest and churn fastest — what contract structure would hold them?

A: They adopt fast because mobile solves their downtime problem instantly; churn fast because procurement re-bids annually and loyalty is thin. Hold them with SLA-based contracts (response-time guarantees, preventive schedules, reporting) that create switching costs and documented value — not rate loyalty. [§ 11.2 contract design, § 1.1 OrderWinner, § 6.1]

Q8. Customers love the convenience, then book their next service at a shop. What does repeat-rate data say we're selling?

A: A trial: they used you once for the emergency/novelty, then defaulted to the shop habit (or a shop captured them via a declined job). The repeat data says convenience alone doesn't retain — you need the follow-up system (service reminders, maintenance plans) that shops run and you skipped. [§ 2.3 retention mechanics, § 12 CRM, § 13 habit]

Q9. When we can't fix it on-site, why does the tow-and-refer moment lose the customer — whose shop do we hand them to?

A: Whoever answers your phone that day — an unmanaged handoff to a random shop that then owns the relationship. Formalize: one partner shop, a named contact, a warm handoff ("Sarah at XYZ expects you"), and reciprocal terms. The handoff is a designed [Activity](#) or a customer leak. [D5, § 8.2 blueprint, § 11.2 partnerships]

Q10. Does pricing reflect route density or premium convenience — which model breaks first at \$5 gas?

A: If pricing is convenience-premium-based, gas barely matters (fuel is small share); if it's density-thin pricing, \$5 gas destroys the sprawl zones. The model breaks where drive cost meets thin tickets. Stress-test by zip: the outlying zones need minimums, zone fees, or abandonment. [§ 11.4 cost shock, § 10 density, § 2.3]

61. Hair salons & barbershops

Q1. When we raise prices, best clients stay and newest leave — is that the outcome we want?

A: Usually yes: newest clients are the least invested (price-shoppers still auditioning you), best clients buy the relationship. Price increases are a client-portfolio filter. The caveat: if today's new clients are tomorrow's best, too-frequent increases starve the pipeline. Track cohort quality, not just churn count. [§ 2.1 segmentation, § 2.3 pricing as filter, § 1.3 BalancedScorecard customer perspective]

Q2. Why does the retail shelf carry products stylists like rather than what clients buy online afterward?

A: Because buying decisions were delegated to staff preference, not demand data. Pull online-after-visit purchase data (ask, or sample client mentions): stock what clients actually use, and align stylist incentives with selling those. The shelf is inventory; it should be demand-driven, not taste-driven. [§ 5.1 inventory, § 13 incentives, § 2.2 demand data]

Q3. Read the lease and commission splits together: service business renting chairs, or brand renting talent?

A: The documents will confess: high commission + client lists owned by stylists = you're a landlord to talent; brand-owned clients + training + walk-in flow = a real service business. Both work; the failure mode is paying brand-level costs while owning landlord-level assets. [§ 1.1 StructuralDecisions, § 13, § 4.1 intangible assets]

Q4. Saturday demand exceeds capacity while Tuesday idles — what have we done that moves demand rather than discounting Tuesday?

A: Probably nothing structural: discounts recruit price-sensitive new clients to Tuesday while regulars keep Saturday. Real demand movers: shift regulars' booking habits (standing-appointment incentives for off-peak), tiered pricing by slot, and stylist schedules matched to the peak. [§ 2.3 demand shaping, § 8.2 yield/fences, § 7 scheduling]

Q5. Clients follow departing stylists at 60%. If the brand retains nothing the chair doesn't, what is the brand?

A: A landlord with a logo. Brand value must live in what stylists can't take: booking system convenience, walk-in demand, guarantee policies, retail ecosystem, training. Build institutional touchpoints (front-desk relationships, reminder systems, loyalty points) so clients have two relationships, not one. [§ 13 relationship capital, § 1.1 structural, § 11.4 key-person risk]

Q6. When a stylist takes maternity leave, half her clients stay in-house, half vanish. Which half did we plan for?

A: Neither, apparently — the plan should be: pre-leave warm handoffs (client meets the covering stylist with the original present), preference notes in the booking system, and a return-date communication cadence. The vanishing half never got a bridge. [§ 8.2 blueprint/handoff design, § 13 knowledge transfer, § 2.3]

Q7. Does online booking fill the book or shift no-show risk to strangers we've never spoken to?

A: Check no-show rate by booking channel: online-first-time clients typically no-show at multiples of phone/returning clients. Counter with new-client deposits or confirmation calls for online-first bookings — convenience for known clients, verification for unknown ones. [§ 8.1 abandonment, § 8.2 fences, § 13 commitment]

Q8. Why do color clients rebook at checkout while cut clients say "I'll call" — which habit does the front desk work?

A: Color clients have a visible regrowth deadline (the roots enforce the calendar); cuts degrade gradually with no urgency marker. The desk works whoever's easy. Fix: give cuts a deadline frame ("your shape lasts about 5 weeks — let's hold your spot") and prebook both equally. [§ 13 urgency/framing, § 2.3, § 8.2 blueprint]

Q9. Junior stylists' books fill with price-sensitive clients who leave when the junior earns a raise. Is the program building or borrowing a clientele?

A: Borrowing: you routed price-shoppers to the junior's lower tier, and the raise reprices those clients out — they leave to the next junior. Building requires routing clients by fit and growth potential, with a designed transition story when prices step up ("now with advanced certification"). [§ 2.3 price-tier architecture, § 13, § 1.3]

Q10. Why do booth-renters outsell commission stylists in retail — the thing commission was designed to drive?

A: Because renters are owners: every retail dollar is theirs, so they sell; commission stylists get a sliver and rationally prioritize chair time. The incentive teaches the behavior. If retail matters, make commission on retail meaningful or gamified; if not, stop pretending the model drives it. [§ 13 incentives/Goodhart, § 1.1 infrastructural decisions, § 2.3]

62. Nail salons

Q1. Why does tip income vary 40% while quality scores are flat — what are clients tipping for?

A: The relationship layer: conversation, memory of preferences, the feeling of being a regular — not the nails (flat quality scores). Tips price the interpersonal product. This means your retention strategy should protect relationship continuity (same tech per client) more than technical polish.

[§ 6.3 SERVQUAL empathy, § 13, § 2.3]

Q2. Why do high-skill artists draw from 20 miles while basic manicures lose to the salon next door?

A: Different products: art is a destination purchase (skill-differentiated, worth travel); the basic mani is a convenience commodity (nearest wins). You can't win the commodity war on quality the customer can't judge — win it on price/speed/location, or exit it and be the destination studio.

[§ 1.1 OrderWinner by product, § 10 location, § 6.1]

Q3. When we matched the new competitor's prices, we lost the bottom third of our book, not the fringe. Why?

A: Because the "bottom third" wasn't price-sensitive — they were value-anchored: your lower price was part of their perception of smart choice, and matching the newcomer signaled you'd been cheap, not good. Price moves reposition you; know what your price meant before moving it. [§ 13

anchoring/signaling, § 2.3, § 6.1 perceived quality]

Q4. Punch-card holders and non-holders visit at identical frequency. What is the card doing — and what would we do with its margin?

A: The card is subsidizing existing behavior (breakage-free discount to regulars who'd come anyway). If frequency is unchanged, kill the card and reinvest its margin in something that changes behavior: a bring-a-friend program or reactivation offers to lapsed clients. [§ 13 Goodhart, § 2.3 loyalty

design, § 1.3]

Q5. Why does December (busiest month) generate the worst Yelp reviews — capacity or seasonal clients?

A: Both: capacity strain (longer waits, rushed work — Kingman at the chair) plus one-time clients (gift bookings, event nails) with no relationship patience. Protect December: cap bookings below max, reserve slots for regulars, and staff for the review you want in January. [A4, § 8.2

yield/segmentation, § 6.3]

Q6. Clients of a departed tech try exactly one other tech before leaving the salon. What happens in that visit — who owns it?

A: An audition nobody manages: the client compares silently, the substitute tech doesn't know the preferences, and the salon treats it as a routine booking. Own it: brief the substitute (notes, colors, habits), comp a small upgrade, and follow up after. One managed visit saves the client; one default visit loses her. [§ 8.2 service recovery/handoff, § 13 relationship capital, § 12 CRM notes]

Q7. Why do walk-ins get served in minutes on weekdays while appointment holders wait — which customer does the floor prioritize?

A: The visible one: walk-ins wait in the lobby where their impatience is watched; appointment holders' time is invisible. You've inverted the priority — appointments are your committed demand and deserve the protected slots. Enforce appointment-first dispatch, even at walk-in cost. [§ 8.1 queue discipline, § 1.1 OrderWinner dependability, § 13]

Q8. Ask ten clients what they paid for: nail care or 45 minutes of sitting still. Which answer does pricing survive?

A: If they say "sitting still" (me-time, sanctuary), your pricing survives anything — you're selling a ritual space, and the nails are the excuse. If they say "nail care," you're in a commodity technical market. The answer tells you whether to invest in ambiance or efficiency. [§ 6.1 perceived quality, § 1.1 positioning, § 8.2]

Q9. Regulars arrive every two weeks like clockwork and never buy retail. What is checkout for?

A: Currently: payment processing. Retail to regulars requires a different mechanism than shelf display — they're on autopilot. Use the service moment (tech demonstrates the cuticle oil mid-service) and make checkout confirm the next appointment; the rebook IS your retail. [§ 2.3, § 8.2 blueprint, § 13 habit]

Q10. Why do dip-powder clients stay twice as long per visit but rebook at the same interval — compressing chairs-per-day?

A: Because dip lasts longer (fewer visits/year) AND takes longer (more minutes/visit) — the service mix shift quietly halves chair productivity while revenue per visit rises less than proportionally. Price dip to its true chair-time cost, or push express services into the mix it displaces. [§ 4.2 capacity/time accounting, § 1.3, § 2.3 pricing]

63. Day spas & massage therapy

Q1. Why does the intake form collect health data nobody reads while contraindication incidents keep happening?

A: Because the form is a liability ritual, not a process input: the data enters no workflow (no alert, no pre-session review). Make contraindications a system flag the therapist must acknowledge before the session — collection without consumption is [NVA](#) that also fails to protect. [D3 NVA, § 12 workflow integration, § 14 safety, § 6.3 poka-yoke]

Q2. Why do couples bookings spike weekend revenue but produce almost zero individual repeats?

A: Because couples massage is an occasion product (anniversary, gift) — the unit of demand is the event, not the person. Convert by capturing each individual (separate follow-ups, solo-visit offers), not the couple. Otherwise accept it as occasion revenue and price it as such. [§ 2.1 demand type, § 2.3 conversion design, § 13]

Q3. Why do massage memberships retain while facial memberships churn with identical discounts?

A: The discount is identical; the product rhythm isn't: massage delivers felt relief monthly (renewal of a felt need), facials promise cumulative skin outcomes (faith-based, delayed). Churn follows the visibility of benefit. Fix facials with progress tracking (photos, skin metrics) so benefit becomes felt.

[§ 6.1 perceived benefit, § 13, § 8.2 subscription design]

Q4. Why do gift-card recipients convert to regulars under 15% despite full-quality experience?

A: Because the gift frame makes it an event, not a habit entry: the recipient didn't choose the category, paid nothing, and feels no pull to return. Convert at the visit: first-visit membership pricing offered at checkout ("today's visit becomes free if..."). The frame breaks at the first personal payment. [§ 13 behavioral/framing, § 2.3, § 8.2]

Q5. Does relaxation positioning compete with medical massage for the same client — which does booking screen for?

A: It doesn't screen, so both clients get the generic booking: the pain client ends up with a relaxation therapist (bad outcome, no return) and vice versa. Add a triage question at booking (relaxation vs. problem-focused) and route to matching therapists — mismatch is your silent churner. [§ 8.2 blueprint/triage, § 6.3, § 4.4 skill routing]

Q6. What fills Tuesday 2 PM: wellness-seekers or occasion-buyers — which has marketing ever targeted?

A: Neither — marketing sells occasions (weekends, gifts) while off-peak needs wellness regulars (retirees, shift workers, chronic-pain clients). Off-peak yield strategy: membership tiers with weekday-only pricing and targeted outreach to schedule-flexible segments. [§ 8.2 yield/off-peak, § 2.3, § 2.1 segmentation]

Q7. Sauna/amenity utilization under 20% while membership sales cite it constantly. What gives?

A: The amenity is a sales asset, not a usage product — it closes memberships by existing, not by being used. That's fine (low utilization = low cost) until members notice they're paying for an unused promise. Bundle usage nudges (sauna-before-session rituals) so the promise stays real.

[§ 6.1 perceived quality, § 13, § 4.2 utilization]

Q8. Prepaid-series clients finish and vanish; drop-ins book for years. Which should the price list encourage?

A: Drop-in rhythm, apparently — the series front-loads commitment and back-loads nothing (completion ends the story). Replace expiring series with rolling memberships (monthly session, pause-anytime): the commitment device without the cliff. Price the list so the membership beats both. [§ 8.2 subscription design, § 13 goal completion, § 2.3]

Q9. Why do rebooking rates range 30–80% with identical training — what does the 80% therapist say in the last two minutes?

A: A specific, dated next step tied to the body: "Your right hip needs two more sessions about three weeks apart — I have the 14th or 16th." The vague "book when you need me" therapists get 30%. Script the last two minutes; it's the highest-leverage process step you own. [§ 13 standard work, § 8.2 blueprint, § 2.3]

Q10. When a therapist builds a waitlist, why do we raise their prices last — what does that cost in chair-time revenue?

A: Because raising their price feels like risking your best asset — meanwhile the waitlist IS the price signal ignored: excess demand at current price. Every month of delay donates the premium. Raise waitlisted therapists' prices first (demand-verified), not last. [§ 8.2 yield/dynamic pricing, § 2.3, § 13 loss aversion]

64. Tattoo studios

Q1. Does the commission split assume a walk-in business with artists, or an artist collective with a storefront — which do artists believe?

A: Artists believe collective (their following, their clients); your split likely prices walk-in traffic you no longer generate. The mismatch surfaces at every renegotiation. Align: if artists bring the demand, splits should approach booth-rent logic; if the shop does, defend it with marketing data. [§ 1.1 StructuralDecisions, § 13 incentives, § 2.1 demand source]

Q2. Flash days fill the shop with first-timers who never return; regulars book six months out. What are flash days for?

A: Cash-flow spikes and apprentice reps — not client acquisition. If flash is marketed as growth, it's failing; if it's capacity-filling on dead days, it's working. Judge it on margin-per-day and artist development, not new-client counts. [§ 2.3 demand shaping, § 4.2 utilization, § 13]

Q3. Does the shop minimum filter bad clients or just small tattoos — what does no-show by deposit size show?

A: Check it: if no-shows concentrate below the minimum, the filter works (small-commitment clients flake); if not, the minimum just refuses small revenue. Deposits are the stronger filter — scale deposit to session size and let the minimum go. [§ 8.1 abandonment, § 2.3 commitment devices, § 13]

Q4. Why do clients negotiate the hourly rate but never the deposit — which protects the artist's time?

A: The deposit — clients intuit the rate is negotiable social terrain while the deposit is a structural fact. That's your lesson: frame the rate as equally structural ("shop rates are fixed; deposit secures the date") and negotiation collapses. [§ 13 framing, § 2.3, § 0.3 governedBy policy]

Q5. After the last artist departure, which revenue stayed: unfinished large pieces or small-piece clients — what does the shop own?

A: Whichever stayed tells you: large-piece clients are artist-bound (multi-session trust), small-piece/walk-in clients are shop-bound. If smalls stayed, the shop owns convenience demand; if nothing stayed, the shop owns walls. Build shop-level assets (brand, flash library, booking system) accordingly. [§ 13 relationship capital, § 11.4 key-person risk, § 1.1]

Q6. Why do touch-up policies get abused by exactly the clients who needed them least — and why never enforced?

A: Because enforcement is socially costly (arguing with a client) and the abusers know it: they picked at the tattoo, claim normal healing, and you fold. Enforcement needs evidence (healed-photo protocol at 2 weeks) and a stated policy at booking — then it's a process, not a confrontation. [§ 0.3 governedBy, § 6.3 acceptance criteria, § 13]

Q7. Why do consult deposits convert at 90% while inquiry-to-consult runs at 10% — what dies in the gap?

A: Undecided browsers meet a slow or generic response: inquiries age unanswered for days, get a price-list reply, and evaporate. The 10% is a response-speed and personalization problem. Answer inquiries within hours with artist-specific portfolios and a low-friction consult booking. [§ 1.1 OrderWinner responsiveness, § 8.2 blueprint, § 12]

Q8. Why do cover-ups — our most demanding work — get quoted at standard rates?

A: Because pricing follows the size/placement table, not the difficulty: cover-ups carry design constraints, rework risk, and consultation overhead the table ignores. Difficulty deserves a multiplier — your scarcest resource (top artists' judgment) is being priced like commodity time. [§ 2.3 value pricing, § 4.4 skill premium, § 1.3]

Q9. When an artist builds an Instagram following, out-of-towners book once while locals sustain the shop. Which does booking policy serve?

A: Design it to serve both deliberately: out-of-town followers fill multi-day guest spots and premium large work (marketing gold), locals fill the base-load. Conflict arises when out-of-town bookings displace local availability and erode the base. Fence guest-spot capacity. [§ 8.2 yield/capacity fencing, § 2.1 segmentation, § 13]

Q10. Why do apprentices' paid practice tattoos generate complaints while portfolios improve — whose education is the client funding?

A: The client's, unknowingly — and that's the complaint: they paid full price for a training rep. Apprentice work needs its own product tier: disclosed, discounted, supervised, with free-touch-up guarantees. Transparency converts the complaint into a bargain-hunter feature. [§ 6.1 perceived quality/honesty, § 2.3 tiering, § 13]

65. Gyms & fitness studios

Q1. Why do PT clients cancel precisely at the moment of success?

A: Because the program was goal-framed: success completes the contract in the client's mind ("I did it — done"). Counter with the next-goal conversation before the goal lands (maintenance phase, performance goals) — the trainer who waits for the success conversation gets cancelled at it. [§ 13 goal-gradient, § 8.2 blueprint, § 2.3]

Q2. What do churn interviews cite: access, community, or transformation — which does the sales script sell?

A: Typically churn cites lapsed habit/no-results (transformation failure) while sales sells access (equipment, hours) — the script sells the wrong product. Sell the first 90 days' transformation plan, not the keycard; retention follows early wins. [§ 1.1 OrderWinner, § 13 habit formation, § 6.3 knowledge gap]

Q3. Members joining with a partner stay 2× longer, and sales never engineers that. What would a bring-a-partner default look like?

A: Make the pair the default SKU: "join together" pricing at signup, partner workout booking at onboarding, dual check-ins in week one. The data already proves the mechanism; the process just never operationalized it. [§ 13 social commitment, § 8.2 blueprint, § 2.3]

Q4. Why do corporate-wellness signups never swipe in — who pays for those memberships?

A: The employer (or the employee via payroll deduction on autopilot): these are breakage memberships — pure margin while unused. The risk is program renewal day, when the employer audits utilization. Get ahead: run engagement campaigns for corporate cohorts so renewal data looks human. [§ 8.2 breakage economics, § 13, § 11.2 B2B renewal]

Q5. Does the class schedule reflect demand data or instructor availability — which would members guess?

A: Instructor availability, and members guess right. Audit bookings per slot against schedule: kill the 5-person Tuesday class, expand the waitlisted ones. Instructor convenience is a real constraint, but the schedule should start from demand and negotiate with instructors, not the reverse. [§ 2.1 demand data, § 7 scheduling, A3]

Q6. January joins quit by March at the same rate regardless of onboarding quality. Are we selling fitness or selling January?

A: January — you're monetizing annual resolve, which decays on a fixed curve. If onboarding can't bend the curve, price for it honestly (front-loaded fees, 3-month commitments) and focus retention spend on the 20% whose habit does form (identifiable by week-4 visit frequency). [§ 13 behavioral cycles, § 2.1 segmentation, § 8.2]

Q7. Why do least-active members never cancel while most-active churn — which does the model depend on?

A: The model depends on the dormant majority (breakage): they pay and never come. The active churn from crowding, boredom, or goal completion. Protect both: keep dormants guilt-free (easy pause, no cancellation friction — or you invite the great unsubscribe), and feed the actives novelty. [§ 8.2 breakage, § 13, § 14 compliance & standards]

Q8. Why does 5–7 PM drive the complaints that cost members, while empty hours drive margins?

A: Peak crowding degrades the experience (wait for racks, packed classes) exactly when your best-attending members show up — and members who attend are the ones who notice. Capacity management at peak (class caps, off-peak incentives) protects the members most likely to leave over it. [A4, § 8.2 yield, § 6.3]

Q9. When a popular instructor leaves, classes recover in 8 weeks but cancellations run 6 months. Why?

A: Attendance is behavior (fills fast with a decent sub); cancellation is identity (the member's relationship with your brand was person-mediated, and the leaving is a slow re-evaluation). Bridge it: celebrate the successor publicly, migrate the community rituals, contact the instructor's core members directly. [§ 13 relationship transfer, § 2.2, § 8.2]

Q10. When we add requested equipment, usage concentrates on the old equipment anyway. Why?

A: Requests came from the vocal few; the majority's behavior is habit-locked to familiar machines. New equipment needs onboarding (floor staff intros, challenge programming) or it becomes expensive décor. Request ≠ demand — measure usage projections before the next capital ask.

[§ 13 stated vs revealed preference, § 4.2 utilization, § 2.2]

66. Yoga & pilates studios

Q1. When a beloved teacher's style changes, attendance holds but retail and membership attach fall. Why?

A: Attendance is habit (the Tuesday 6 PM slot survives); attach is enthusiasm (the community feeling that drove extra spending). The style change cooled the emotional layer without breaking the routine — yet. Watch the leading indicator: attach falls first, attendance follows in a quarter or two. [§ 13 behavioral ops, § 2.2 leading indicators, § 6.3]

Q2. Why does sub-policy chaos cost more members than price increases ever have?

A: Because members buy a specific practice relationship and schedule reliability; a cancelled class or mystery sub breaks the ritual they've built their day around. Price is a number; reliability is the product. Pay for a real sub bench and same-teacher consistency guarantees. [§ 6.3 SERVQUAL reliability, § 1.1 OrderWinner, § 8.2]

Q3. Why do reformer clients pay 2.5× yoga rates with half the retention — which product is the future?

A: Check the retention cause: reformer churn often reflects price fatigue and limited class inventory (machine count caps scheduling). Yoga retains on community; reformer monetizes on scarcity. The future is probably both: reformer as premium tier, yoga as community base — don't let one cannibalize the other's room. [§ 8.2 yield, § 2.1, § 10 capacity]

Q4. Does workshop programming deepen community or monetize the same 20 people — what does attendee overlap show?

A: The overlap tells you: high overlap = a paid club for your inner circle (fine, but call it that); low overlap = real community expansion. If it's the same 20, vary format/price/topic to widen the funnel, or accept workshops as retention perks for your core. [§ 2.1 cohort analysis, § 13, § 2.3]

Q5. Why do intro-special attendees convert to class packs, not memberships, against our incentives?

A: Because packs preserve optionality — the intro customer isn't sure they'll sustain the habit, and a pack caps their downside. Your incentives push membership because it's better for you; their choice is rational for them. Bridge: pack-to-membership conversion offers timed to pack exhaustion. [§ 13 risk/optionality, § 8.2 pricing ladder, § 2.3]

Q6. Why do teacher-training graduates open studios within five miles — threat, compliment, or design flaw?

A: Design flaw if unintended: your TT program trains and certifies competitors while giving them your playbook and community visibility. Options: non-compete-lite (distance clauses where enforceable), alumni network with preferential terms, or price TT to include the competitive risk. [§ 11.4, § 13 knowledge leakage, § 1.1 structural]

Q7. Why do we lose members to home-practice apps at exactly the two-year mark, when practice is strongest?

A: Because competence kills dependence: at two years they have a self-practice vocabulary, and the app is cheaper for someone who no longer needs instruction. Retain the proficient with what apps can't do: advanced workshops, community leadership roles, hands-on teaching. [§ 2.1 lifecycle, § 13, § 1.1 value evolution]

Q8. Why do 6 AM classes fill with highest-retention members while prime evening classes churn — which slot does the schedule protect?

A: 6 AM members are the committed identity-core (practice before anything can interfere); evening members are fitting fitness into life margins — churn-prone by construction. Protect the morning slots' quality obsessively; use evenings for acquisition and variety. [§ 2.1 segmentation by commitment, § 8.2, § 13]

Q9. Which survives a 20% rent increase: a fitness business with a spiritual aesthetic, or a community with a fitness product?

A: The community: if members' identity and friendships live in your room, price tolerance is high and alternatives feel like betrayal; if you're a fitness vendor, 20% sends them to the gym down the street. The answer determines whether to invest in programming/community or amenities. [§ 1.1 positioning, § 6.1 identity value, § 13]

Q10. Unlimited members attend less over time but renew at the highest rates. What are they buying?

A: The option and the identity: "I am someone with a studio membership" plus guilt-free flexibility. Declining attendance is the breakage you bank. Keep the identity alive (occasional "we miss you" that invites without shaming) and the renewal survives the non-attendance. [§ 8.2 breakage/option value, § 13 identity, § 2.3]

67. Martial arts schools

Q1. When the head instructor teaches less, retention falls school-wide, not just in their classes. Why?

A: Because the head instructor IS the brand's quality signal: their presence on the mat calibrates the whole school's standard, and parents read their absence as decline. Delegate authority visibly (named senior instructors with introduced credentials) before stepping back, or the school deflates with your schedule. [§ 13 brand-as-person, § 6.1 perceived quality, § 11.4 succession]

Q2. Adult programs retain through injuries while kids churn at summer break. Which does facility design serve?

A: Check scheduling and space: most schools are designed for the kids' after-school block (the revenue engine) with adults in leftover slots. Kids churn at schedule disruptions (summer) — the counter is summer programming (camps, flexible attendance) that keeps the rhythm unbroken.

[§ 2.1 seasonality, § 2.3, § 10 facility]

Q3. Why do trial students who bring a friend enroll at 3× — and why isn't that the entire trial process?

A: Because the friend converts a scary evaluation (being judged) into a shared adventure — social proof and commitment in one move. Make "bring a training partner" the default trial offer, with paired enrollment incentives. The mechanism is proven by your own data; the process just hasn't canonized it. [§ 13 social commitment, § 2.3, § 8.2 blueprint]

Q4. What do 4 PM parents believe they bought: martial arts education or after-school care with mats — which does retention depend on?

A: Test by watching renewal conversations: if parents cite behavior and discipline, they bought development (retention depends on visible progress); if they cite schedule coverage, they bought childcare (retention depends on convenience and price). Serve whichever is true — deliberately.

[§ 1.1 OrderWinner, § 6.3 knowledge gap, § 2.1]

Q5. Competition-team families spend most and complain most. Which do we manage to — and what does that cost?

A: They complain because they're invested (visibility-obsessed); they spend because they're invested. Managing to their complaints distorts the school for the silent majority who fund it. Manage to the recreational base's quiet satisfaction; serve comp families' intensity without letting it set school policy. [§ 13 vocal-minority bias, § 2.1 segmentation, § 1.1]

Q6. Does black-belt-club pricing create commitment or a dropout cliff at the contract boundary — what does month-37 show?

A: Month-37 (right after typical 36-month contracts) tells you: a cliff means the contract was the commitment (artificial), not the practice. Smooth it with rank-based milestones beyond the contract horizon (the black belt itself takes longer than the contract) so the journey outlasts the paperwork. [§ 13 commitment design, § 8.2 subscription, § 2.3]

Q7. When parents enroll kids for discipline, why do the families who need it most cancel first?

A: Because the discipline-seeking families often have chaotic logistics — inconsistent attendance, payment friction, schedule instability. The kids who'd benefit most attend least. The intervention is operational: attendance tracking with early outreach, carpool networks, scholarship structures that buy consistency. [§ 13, § 2.1 adverse selection, § 8.2]

Q8. Why market anti-bullying to parents while retention data says kids stay for friends?

A: Because acquisition and retention are different conversations: parents buy the anti-bullying promise (their anxiety), kids stay for the tribe (their joy). Both messages are true for their audience — but your programming must deliver the tribe (partner drills, team events) or the parent's promise churns at month 4. [§ 1.1 two-audience positioning, § 13, § 6.3]

Q9. Students quit at belt-test boundaries, not between them. What are we selling between tests?

A: Apparently not enough: if quitting clusters at achieved ranks, the belt is the product and the training is the toll. Make the between-test time valuable: skill curriculum visible weekly, stripes/micro-progress, sparring privileges — progress must be continuous, not ceremony-gated.

[§ 13 goal-gradient design, § 8.2 blueprint, § 6.1]

Q10. Why does retention improve exactly when we stop measuring it, in every new program?

A: Suspicious pattern — likely measurement artifact: new programs get launch attention and a retention dashboard; when attention moves on, the dashboard dies and "retention" is declared from vibes. Or programs improve once they stop being pilot-observed. Either way: make retention measurement permanent infrastructure, not launch theater. [§ 13 Goodhart/Hawthorne, § 0.3 measuredBy, § 6.3]

68. Dance studios

Q1. Why do summer intensives enroll existing students at full price while beginner camps sit half empty?

A: Because intensives sell to proven demand (your own committed families) while camps require external acquisition marketing you don't do — camp buyers are new families who've never heard of you. Camps need a real acquisition channel (schools, community lists) or accept them as conversion-loss leaders priced accordingly. [§ 2.3 acquisition vs retention demand, § 1.3, § 13]

Q2. To the parent writing the check, what justifies pricing: dance education or childhood participation — which would they cancel first?

A: Participation (the costume, the recital, the photos) is nearly uncancellable; "education" is abstract. Your pricing survives on the experience product. That's not cynical — it tells you to invest in the visible milestones families actually renew for. [§ 6.1 perceived quality, § 1.1 OrderWinner, § 13]

Q3. Why do a departing teacher's competition students follow her out while recreational students stay?

A: Comp families bought the coach (intensive, personal, ambition-bound); recreational families bought the program (schedule, friends, recital). Your comp program's loyalty is person-bound — protect with institutional framing (team identity, multiple coaches) or accept that comp revenue walks. [§ 13 relationship capital, § 11.4 key-person risk, § 2.1]

Q4. Does adult-program revenue stabilize the business or distract it — what do scheduling conflicts cost in kids' enrollment?

A: Compute the shadow price: if adult classes occupy prime after-school slots, each adult dollar may cost more in displaced kids' enrollment (your annuity demand). Adults in off-peak slots = stabilization; adults in prime slots = cannibalization. [A2 opportunity cost, § 7 scheduling, § 1.3]

Q5. Why do we lose families to competitors at recital time — our peak showcase moment?

A: Because recital is when every deficiency is visible side-by-side: venue, organization, costume costs, and the child's moment. Families comparison-shop during your showcase. Audit the recital as a retention product (smooth logistics, fair costs, every child featured) — it's your biggest renewal event, not just a show. [§ 6.1 perceived quality moments, § 8.2 blueprint, § 13]

Q6. Why do costume/recital fees — most resented charges — coincide with highest retention season?

A: Because the fees purchase the emotional peak (the recital) that locks the next year's commitment — resentment and retention coexist because the value is real but the billing is clumsy. Bundle fees into tuition visibly ("all-inclusive") and the resentment dissolves while retention stays. [§ 13 framing/mental accounting, § 2.3 price architecture, § 8.2]

Q7. Why does the waitlist for popular classes coexist with empty slots one level up — what does placement get wrong?

A: Promotion criteria are too rigid or too slow: students ready to advance wait for annual evaluations while the popular level bulges. Continuous assessment with monthly promotion windows unblocks the ladder — the waitlist is a process defect, not just demand. [§ 7 flow/blockage, § 3.4 Blocked state, § 13]

Q8. Why do the most talented students' families negotiate hardest on tuition — what does the scholarship budget buy?

A: Because talent opens options (other studios court them) and they know their child is marketing value (competition trophies, recital solos). The scholarship budget buys your studio's visible excellence — legitimate, but cap it and tie it to ambassadorial duties. [§ 13 negotiation dynamics, § 2.3, § 1.1]

Q9. Recreational dancers fill classes while competitive fees carry the studio. Which identity does the recital serve?

A: Watch the recital: if it showcases comp numbers and trophy routines while recreational kids get filler slots, it serves the minority revenue and alienates the majority base. The recital is your retention engine for the recreational majority — program it for them. [§ 1.1 identity/focus, § 6.1, § 13]

Q10. Why do age-4–6 starters stay a decade while age-10 starters quit within a year — where does marketing spend go?

A: Early starters integrate dance into identity before social competition and self-consciousness peak; 10-year-old starters compare themselves to 8-year veterans and quit. Market heavily to preschool families (the decade annuity); for older starters, sell beginner-cohort classes (no comparison gap) instead of level placement. [§ 13 identity formation, § 2.1 lifecycle, § 2.3]

69. Real estate brokerages**Q1. Top producers generate our lowest per-agent margins. Which does the split structure reward?**

A: Production, at any cost: the high-split tiers that attract top producers concede your margin to win their volume — rational if their brand gravity recruits others, corrosive if it just buys pride. Calculate producer value including recruitment halo; if the halo is myth, renegotiate at renewal.

[§ 13 incentives, § 1.3, § 1.1 structural]

Q2. When the market slows, part-time agents outlast full-timers. Who is this model built for?

A: Part-timers: your fee structure (low fixed, high split) subsidizes agents with other income, while full-timers carry the desk fees and die in downturns. That's a valid model (volume of licenses) — but don't confuse license count with a productive salesforce. [§ 2.1 cycle exposure, § 1.1 business model, § 13]

Q3. Listings sell fastest in the price band where we have least agent experience. What is the market telling us?

A: Where to recruit: the fast band (likely entry-level homes) is where transaction velocity and buyer demand concentrate — your experienced-agent-heavy roster is positioned for a slower, prestige market. Recruit hustlers for the velocity band or train existing agents down-market. [§ 2.2 demand signals, § 1.1 positioning, § 13]

Q4. When a team forms inside the brokerage, why does our margin on them fall while their loyalty falls with it?

A: Because teams internalize the value you provided (leads, admin, brand) — they renegotiate splits from strength and identify as the team, not the brokerage. Teams are your successor-competitors in embryo. Price team services explicitly (per-seat fees, lead charges) and give them reasons to stay that scale with their size. [§ 11.2 internal make-vs-buy, § 13, § 1.1]

Q5. Why do agents join for the brand and leave for the split — at what production level does math flip?

A: The flip is where brand-derived leads fall below the split delta: once an agent's book is self-generated (referral-based), the brand contributes nothing visible and the split is pure tax. Compute it per cohort; retention above the flip requires non-split value (leads, staff support, office). [§ 2.3, § 13 incentives, § 1.3]

Q6. Why do we recruit with promises we measure and retain with culture we don't?

A: Because recruitment is a transaction (measurable: splits, leads) and retention is a relationship (unmeasured: belonging, support). The unmeasured side is where they leave through. Instrument culture: manager touchpoint cadence, agent satisfaction pulses, support-ticket response. [§ 0.3 measuredBy, § 13, § 6.3 SERVQUAL internally]

Q7. Why do signed buyer agreements still defect at the listing-appointment stage?

A: Because the agreement is paper while the relationship is vibes: when they list, the listing agent's sphere and the seller-agents they interview re-open the loyalty question. Buyer reps need ongoing value delivery (market updates, off-market access) so the listing decision starts with you, not an open audition. [§ 13 relationship maintenance, § 2.3, § 6.3]

Q8. Why do sphere-based agents survive every market while lead-gen agents churn every downturn — which do our fees assume?

A: Sphere agents own a renewable asset (relationships); lead-gen agents rent demand from you or portals — and rent stops being payable in downturns. Your fee model likely assumes the renter (high split covers lead cost). Long-term, build sphere development (referral systems) into agent training. [§ 2.1 demand ownership, § 11.4 cycle resilience, § 13]

Q9. Read the marketing budget as strategy: does it fund agent services or a consumer brand?

A: Whichever it funds is your real model: consumer-brand spend only pays if it generates attributable leads agents can feel; agent-services spend (signs, photography, coaching) retains producers directly. Most brokerages fund a consumer brand agents don't credit and underfund the services they'd stay for. [§ 1.1 deliberate strategy, § 13 attribution, § 1.3]

Q10. Does training produce agents who stay or agents who get recruited — what does 24-month cohort retention show?

A: The cohort data tells you which business you're in. If trained agents get poached at month 18, your training is the industry's free academy — respond with golden handcuffs (deferred bonuses, mentorship equity) or accept academy economics (charge for training, recruit the next cohort).

[§ 13, § 11.4 talent risk, § 1.3]

70. Property management

Q1. Does screening rigor reduce evictions or just extend vacancy — what does all-in cost per placed tenant by criteria tier show?

A: The tier analysis is the answer: strict criteria cut eviction cost but add vacancy days; loose criteria do the reverse. The optimum is where marginal eviction-risk reduction equals marginal vacancy cost — computable from your own placement history. Most managers set criteria by fear, not math. [§ 5.2 Newsvendor logic, § 11.4 risk-cost trade-off, § 1.3]

Q2. Why does leasing speed beat the market while renewal rates lag it?

A: Because leasing is a sales process (optimized, measured, bonused) while retention is an operations process (maintenance response, communication) that nobody bonuses. Fast leasing also admits marginal tenants your operations then lose. Balance the incentives; renewals are cheaper than turns. [§ 13 incentives/Goodhart, § 1.3 turn cost, § 6.3]

Q3. Why do tenants rate us worst at move-out — the moment we stop serving them?

A: Because move-out is where money changes hands adversarially (deposit deductions) with no service recovery afterward — the relationship ends on a dispute. The review lives forever; the tenant is gone. Fix the process: pre-inspection walkthroughs, itemized deductions with photos, fast deposit returns. [§ 8.2 service recovery, § 6.1, § 13 peak-end rule]

Q4. Why do PM portfolios grow past response-quality collapse — what number have we refused to enforce?

A: The doors-per-PM cap: everyone in the industry knows the range (150–250 doors) where quality holds, but adding doors to an existing PM is free revenue while hiring is a cost. Enforce the cap as policy or accept that your growth plan is a service-degradation plan. [§ 4.2 capacity limits, A4, § 0.3 governedBy policy]

Q5. Why do takeovers of mismanaged properties spike first-year costs predictably while proposals assume steady-state?

A: Because you're quoting the destination, not the journey: deferred maintenance, tenant-quality resets, and record cleanup are the takeover's true first-year content. Price takeovers as a distinct product (onboarding fee + remediation budget) — the predictable spike becomes revenue instead of margin erosion. [§ 2.3 product design, § 11.4 inherited risk, § 1.3]

Q6. Which is billable: managing properties or managing owner anxiety — which consumes the hours?

A: Anxiety consumes the hours (calls, explanations, reassurance) while the contract bills for property tasks. Make the invisible visible: monthly reports that preempt the anxiety calls, and communication SLAs — or price a communication tier. Unbilled emotional labor is still a cost. [D3 VA vs NNVA, § 6.3 assurance, § 2.3]

Q7. Why does our maintenance markup — most criticized fee — fund the responsiveness owners value most?

A: Because the markup subsidizes the coordination infrastructure (vendor network, 24/7 dispatch) that the monthly fee can't carry. Owners see the markup line and not the system it funds. Show the math at renewal: $\text{markup} \div \text{responsiveness metrics} = \text{the deal they're actually getting}$. [§ 13 framing, § 2.3 price architecture, § 6.1]

Q8. Why do owners leave for self-management after we've fixed their property — exactly when the job got easy?

A: Because your visible work disappears when the property stabilizes: no crises, no visible effort, and the monthly fee looks like a tax on a quiet asset. Counter with stewardship reporting (what we monitored, prevented, optimized) — make stability legible as your product. [§ 6.1 perceived quality of prevention, § 2.3, § 13]

Q9. Why do largest owners get worst per-door margins and best service — which do they tell their peers about?

A: They get volume discounts plus disproportionate attention (you can't afford to lose them) — and they tell peers about the service, while your margins tell you about the discount. Portfolio pricing should include service-level tiers: if they want white-glove, price it; don't donate it. [§ 2.3 tiered pricing, § 1.3, § 13]

Q10. Silent owners renew; engaged owners who love our responsiveness leave. Which behavior does the model reward?

A: It rewards neglect, perversely: silent owners never test you; engaged owners eventually hit a failure (one bad vendor job, one missed call) and their high expectations convert disappointment into exit. The model needs expectation calibration for engaged owners — over-communication on failures buys more than success streaks. [§ 8.2 recovery, § 6.3 expectation management, § 13]

71. Insurance agencies

Q1. Why do claims-advocacy moments — what clients say they pay for — happen so rarely that no client experiences one before leaving?

A: Because claims are rare events (a client files every few years), so your core differentiator is experientially invisible at renewal time. Manufacture proxy evidence: annual reviews that quantify "coverage gaps we closed, discounts we found" — make advocacy legible between claims, or clients renew on price alone. [§ 6.1 perceived quality of rare events, § 2.3, § 13]

Q2. Does the account-review process prevent churn or just document it — what does save rate by review timing show?

A: Timing data tells: reviews triggered after the renewal notice arrives are autopsies; reviews at 90 days pre-renewal are retention tools. If your calendar keys reviews to the renewal date, you've built a churn-documentation system. Move the trigger. [§ 2.2 leading vs lagging, § 8.2 blueprint, § 13]

Q3. Producers' year-one books look identical; year-five books diverge 5×. What happens in year two — why is it unmanaged?

A: Year two is when launch support ends (house accounts, mentor attention) and the producer's own prospecting system — or absence of one — takes over. It's unmanaged because management watches new-producer ramp, not post-ramp decay. Institute year-two pipeline audits and activity standards. [A9 learning curve completion, § 13, § 2.2 leading indicators]

Q4. Why do commercial clients refer constantly while personal-lines clients — served more often — refer almost never?

A: Commercial buyers experience you as risk counsel (advice = talkable); personal-lines buyers experience a transaction (renewal = invisible). Personal-lines referrals need engineered moments: claims follow-ups, life-event triggers (new home, teen driver) with a referral ask attached. [§ 2.3 referral mechanics, § 6.1, § 13]

Q5. When we win an account on price, why does it underperform on retention even with flat pricing?

A: Selection: price-won clients are shoppers by disposition — they left their last agent on price and will leave you the same way; their carrier also re-rates them upward at renewal, re-triggering the shop. Price-won business needs immediate value-layering (reviews, bundles) to convert shoppers into relationships. [§ 2.1 adverse selection, § 13, § 8.2]

Q6. When a carrier raises rates, why do we lose our best-loss-ratio clients first — who makes the leave decision?

A: Your best clients have the most options (clean records = every carrier wants them), so they can act on rate anger; your worst risks are stuck. The leave decision is made by the client, but the save is made by you in the 60 days pre-renewal — proactively re-market your clean accounts before they shop themselves. [§ 2.1 selection dynamics, § 11.3 carrier behavior, § 2.3]

Q7. Does our comp structure describe a carrier distribution channel or a client advisor — whose interests does it pay us to serve?

A: Follow the money: contingent commissions and carrier bonuses pay for volume and retention with the carrier, not fit for the client. You're a distribution channel with advisory branding. That's the industry norm — but know it, because fee-for-service advisors are eating the high-trust segment. [§ 13 principal-agent, § 11.3 incentive alignment, § 1.1]

Q8. Why does the book grow via agency acquisitions whose clients then churn at 2× organic?

A: Because acquired clients chose the seller, not you: the transition letter is not a relationship, and rate changes at first renewal hand them a shopping trigger. Acquired books need a retention program (personal outreach, coverage reviews in year one) priced into the acquisition multiple. [§ 11.2 M&A integration, § 13 relationship transfer, § 2.2]

Q9. Why does cross-selling life/benefits work when the producer does it and fail when the specialist does?

A: Trust transfer doesn't survive the handoff: the client's relationship is with the producer, and the specialist arrives as a stranger with a product. Have the producer open and stay in the conversation (three-way meetings), positioning the specialist as their team's expert — not a referral-out. [§ 13 trust transfer, § 8.2 blueprint, § 11.2]

Q10. Highest-premium accounts shop us every renewal; smallest never leave. Which does the service model justify?

A: Neither extreme is right: whale accounts shop because your service is undifferentiated at their scale (they need risk-management depth to stay), and tiny accounts stay from inertia (cheap but unprofitable to serve). Segment service: risk-advisory for whales, efficient automation for the tail. [§ 2.1 segmentation, § 8.2 yield, § 1.1]

72. Mortgage brokerages

Q1. Why do loan officers' pipelines collapse when they vacation — what system was never built?

A: A coverage system: pipelines are person-bound (realtor texts, borrower calls all route to the LO's phone). Build pod coverage (paired LOs, shared pipeline visibility, handoff protocols) — a pipeline that dies when its owner rests is a single-point-of-failure org design. [§ 11.4 redundancy, § 13, § 12 CRM workflows]

Q2. What survives the next rate environment: our advice or our access to capital — which are we compensated for?

A: Access is compensated (lender-paid), advice is not — yet advice is what survives rate cycles (refi booms make access a commodity; purchase markets and rate-lock anxiety make advice scarce). Build the advisory identity (and where possible, fee components) before the cycle forces it. [§ 1.1 structural positioning, § 2.1 cycles, § 2.3]

Q3. When we lose deals to retail banks, why is it never on rate — what do borrowers cite?

A: Certainty and simplicity: "my bank had my accounts, it was just easier." You lose to friction, not price. Counter with process transparency (milestones, single point of contact, doc-upload simplicity) — your rate advantage is invisible if the process feels risky. [§ 6.3 SERVQUAL assurance, § 8.2 blueprint, § 13]

Q4. Realtor partners send their hardest files and easiest gratitude. Which one pays?

A: Hard files pay if they close (they cement the relationship — you become the rescue call) but destroy economics if they don't (worked hours, zero fee). Track close rate by file difficulty per realtor; renegotiate the mix with realtors who send only junk: easy files earn the rescue service. [§ 11.2 channel economics, § 1.3, § 13]

Q5. Why do our pre-approvals close at 55% vs competitors' 70% — where in the 30 days do we lose them?

A: Map the fallout by stage: usually it's days 1–7 (borrower still shopping — your pre-approval was a free option) and the appraisal/conditions phase (communication gaps kill confidence). Tighten day-one lock-in conversations and milestone communication; measure fallout by week. [§ 8.2 blueprint, § 2.2 funnel analysis, § 6.3]

Q6. Does our speed-to-close come from process or from declined files — would our realtors agree?

A: Ask them — they know: if your speed stats exclude the files you turned away, the realtors carrying those rejects experience you as slow AND picky. Speed built on selection is marketing, not capability. Publish speed by file tier; realtors respect honest tiering. [§ 13 metric gaming, § 6.1 honesty, § 1.3]

Q7. Why does cost-per-loan from purchased leads exceed realtor-relationship cost 4× — which does the budget grow?

A: Usually purchased leads (they scale with a credit card; relationships scale with time). The 4× gap says relationship development is underfunded: dedicated realtor-liaison effort, co-marketing, and speed-for-partners. Fund the channel with the better economics the patience it requires. [§ 1.3 CAC by channel, § 2.3, § 13]

Q8. When refis surge, why does the purchase pipeline decay — who maintains it?

A: Nobody: refi demand arrives served on a platter (rate alerts, easy closes), so LO attention migrates and realtor relationships starve. When rates normalize, the purchase pipeline is empty. Ring-fence purchase-business activity (realtor touches per week) as non-negotiable during refi booms. [§ 2.1 demand cycles, § 13 attention allocation, § 7 capacity]

Q9. Why do past-client refi campaigns underperform generic aggregator rate alerts?

A: Timing and trigger precision: aggregators fire on rate movements matched to the borrower's actual note rate; your campaigns fire on your marketing calendar. Build rate-triggered outreach from your loan data (note rate vs. market) — you have better data than the aggregator; you just don't use it. [§ 12 data/automation, § 2.2 leading indicators, § 13]

Q10. Why do five-star borrowers refi with whoever calls next — what did the five stars measure?

A: The closing experience (your service), not loyalty (your future relationship). Five stars measure a completed transaction's pleasure; the next refi goes to whoever is present at the trigger moment. Loyalty requires being there then: annual reviews, rate monitoring, actual contact between loans.

[§ 13 Goodhart, § 2.3 lifecycle, § 6.3]

73. Financial advisory firms

Q1. Why does compliance documentation grow yearly while the client-preferences file stays thin?

A: Because compliance has an auditor and preferences have no owner: the regulator reads one file, nobody reads the other. Yet preferences (goals, family, fears) are what makes advice feel personal. Assign the preferences file the same review cadence as compliance — it's the retention document.

[§ 0.3 governedBy vs measuredBy, § 13, § 6.3 empathy]

Q2. Does quarterly-review cadence deepen relationships or train clients to evaluate us quarterly — what does attrition by frequency show?

A: Check it: over-frequent reviews train performance-evaluation (each quarter is a referendum on returns you don't control); sparse reviews drift into irrelevance. The attrition curve usually favors semi-annual depth plus event-driven contact. Cadence is a design choice, not a virtue. [§ 13

Goodhart/evaluation framing, § 2.3, § 8.2]

Q3. Referrals come from happy clients' adult children, not the clients themselves. Who actually experiences our value?

A: The next generation — they watch you protect their parents and decide you're trustworthy. That makes them your future book AND your referral engine. Serve them deliberately: family meetings, next-gen education, and make sure they're in your CRM as relationships, not footnotes. [§ 13, § 2.1

generational segments, § 11.4 succession of the client base]

Q4. If markets delivered neither returns nor comfort for three years, what would clients say they paid for?

A: Whatever you can name now: discipline (not panic-selling), tax management, planning, access to you in fear moments. If the honest answer is "nothing," your fee is a beta harvest exposed at the next flat market. Build and document the non-return value annually. [§ 6.1 perceived quality, § 1.1

OrderWinner, § 2.3]

Q5. When markets drop 20%, why does inbound call volume predict retention better than outbound?

A: Because inbound calls are fear expressing itself — clients who call and get calmed stay; the silent ones quietly move assets later. Your outbound volume is activity; their inbound is signal. Treat every inbound fear call as a retention save with a follow-up loop. [§ 2.2 leading indicators, § 8.2 recovery, § 13]

Q6. Why do AUM fees look stable while revenue per advisor hour falls — what fills the hours?

A: Service creep: more meetings, more planning updates, more regulatory paperwork per client — fee stays flat, hours inflate. Measure hours per household; the fix is service tiering (meeting frequency by asset level) or team leverage (associate-led service for smaller accounts). [§ 1.3 productivity, § 2.3 tiering, § 4.4]

Q7. Clients consolidate assets during fear and attrit during calm. Which phase does the service calendar cover?

A: Probably fear (reviews cluster in volatility) — but the calm is when they drift to the robo-advisor or the golf-buddy's guy. Cover the calm: proactive planning work (Roth conversions, estate reviews) happens in calm markets and creates the stickiness that survives the next one. [§ 2.1 cycle-aware service design, § 2.3, § 13]

Q8. When a client dies, why do heirs transfer out within 18 months at 70% — whose relationship was it?

A: The deceased's — the heirs know you as "dad's advisor," a legacy contact, not their advisor. The 18-month clock starts at death but the relationship was buildable for years before: beneficiary introductions, family meetings, next-gen accounts. Retention of the estate is earned pre-mortem. [§ 13 relationship continuity, § 11.4, § 2.3]

Q9. Why do youngest advisors serve oldest clients while principals take new money — what does that do to continuity?

A: It inverts succession: your oldest clients (nearest to asset-transfer events) bond with your least powerful advisors, while principals hoard the new relationships. When the client dies or the junior leaves, the account is orphaned twice. Deliberately pair seniors with high-stakes households. [§ 11.4 succession risk, § 13, § 1.1 structural]

Q10. Why do planning-first clients stay through fee increases while investment-first clients shop us yearly?

A: Because planning relationships are multi-dimensional (goals, taxes, estate — no easy comparison) while investment relationships reduce to one comparable number (returns vs. benchmark). Planning-first is the moat; investments-first is the commodity. Convert investment clients by leading every review with planning. [§ 1.1 OrderWinner, § 6.1, § 2.3]

74. Home inspection

Q1. Why do agents refer us until we find a deal-killing defect — what does referral flow by report-severity look like?

A: It probably shows agents routing "must-close" deals to softer inspectors after your thorough report kills one — a selection filter you can't see without the data. The counter: be thorough AND fast/communicative so the cost of honesty to the agent is minimized; the best agents (repeat, professional) value deal-survivable honesty. [§ 13 incentive conflict, § 6.3, § 11.2 channel dynamics]

Q2. Buyers who attend the full inspection negotiate better and complain less — attendance is 40%. What moves it?

A: Scheduling design and agent framing: book inspections at attendance-friendly times (late afternoon), tell buyers the walkthrough is where the value is ("the report is a summary; the walkthrough is the education"), and brief agents to set that expectation. Attendance is an engineered outcome. [§ 8.2 blueprint, § 13, § 2.3]

Q3. When a listing agent pre-inspects, why do findings shrink while fees hold — which party do we serve that day?

A: You serve whoever hired you — and pre-listing inspections create pressure to soften (the seller is your client that day). The professional answer: identical standards regardless of side, stated upfront. If your findings shrink, you've priced your integrity into the fee structure. [§ 13 principal-agent, § 6.3 standards, § 14 compliance & standards]

Q4. Our E&O claims come from items we noted in the report. What do clients and courts actually read?

A: The summary and the photos — not the body where your noting lived. If material findings are buried on page 34, noting ≠ communicating. Restructure: severity-tiered summary, photos inline with each major finding, and a verbal walkthrough of the top items. [§ 6.3 communication, § 0.3 Output spec, § 14 liability]

Q5. Why do add-on services (radon, sewer scope) attach at 4× different rates across inspectors — sales skill or liability comfort?

A: Usually liability comfort disguised as modesty: some inspectors avoid recommending scopes because finding something creates follow-on responsibility. Reframe: add-ons are client protection, and the script is risk-based ("this home's age makes the sewer line a \$200 question worth asking"). [§ 13, § 6.3 assurance, § 2.3]

Q6. Does the 60-page report protect us or bury the signal — which complaints reference the report nobody read?

A: Both: the length is legal armor, the burial is a service failure. Complaints cite "it was in the report but we didn't understand it." Solution: layered output — 2-page decision summary + full detail appendix. Protection lives in the appendix; communication lives in the summary. [§ 6.1, § 0.3 specifiedBy, § 8.2 blueprint]

Q7. Why do we book solid two weeks out in spring while scheduling software sits idle in the off-season we never productized?

A: Because inspection demand follows transaction volume (spring market) and you've built no counter-seasonal product: off-season products exist — pre-listing inspections, home-maintenance inspections, radon season, investor walkthroughs. Productize November–February or accept the seasonality as structural. [§ 2.1 seasonality, § 2.3 counter-seasonal, § 7 capacity]

Q8. Same-day reports improve reviews but not agent referrals. Which customer are we optimizing for?

A: The buyer (reviews), not the agent (referrals) — agents value predictability, deal-safety, and communication more than raw speed. If referrals pay more than reviews, optimize the agent experience: proactive status updates, clean scheduling, calls before the report lands. Know your true customer. [§ 1.1 OrderWinner by channel, § 2.3, § 13]

Q9. What do referring agents believe we sell: risk assessment or transaction lubrication — what happens when we deliver the first too well?

A: Many agents quietly want lubrication: inspections that inform without killing deals. Deliver pure risk assessment "too well" and lubrication-seeking agents defect to softer inspectors. Your choice: be the thorough inspector and cultivate the agents (top producers, ethical ones) who want exactly that. [§ 13 channel incentives, § 1.1 positioning, § 6.3]

Q10. Our price is 20% above competitors; agents send difficult clients to us and easy listings to the cheap guy. What does that mean?

A: You're positioned as the specialist (hard cases justify premium) while volume flows to commodity pricing. Two plays: lean in (expert positioning, luxury/complex niche) or build a standard tier to recapture the easy volume. The current split means the market has already segmented you — choose consciously. [§ 1.1 positioning/segmentation, § 2.3 tiering, § 13]

75. Title & escrow companies

Q1. What does office location and layout advertise: compliance utility or relationship business — which do agents choose on?

A: Agents choose on the officer relationship and closing experience — your office either supports that (comfortable signing rooms, hospitality) or undermines it (DMV aesthetics). Agents choose on relationship; the office is the relationship's stage. Audit what your lobby communicates. [§ 6.1 tangibles (SERVQUAL), § 10 facility, § 1.1]

Q2. When a closing goes sideways, agent blames us, lender blames agent — who owns the timeline in the client's eyes?

A: Whoever communicated least — the client assigns ownership to the silent party. Proactive timeline updates (even "no news" updates) make you the competent party by default. In multi-party processes, communication volume IS perceived ownership. [§ 6.3 responsiveness, § 8.2 blueprint, § 13]

Q3. Builder/developer accounts — highest volume — get our most junior closers. What would an audit reveal?

A: That assignment logic follows seniority-perceived-prestige backwards: builders seem "easy" (repeat volume, same forms) until a plat issue or construction-draw complexity explodes under a junior. Assign by complexity-weighted risk, not by glamour. [§ 4.4 skill routing, § 11.4 risk, § 13]

Q4. What happens to the book six months after a senior escrow officer leaves, given agents choose the officer?

A: The book walks — agents follow officers because the officer IS the product (competence, calm, problem-solving). Institutional defenses: team-based service (agents know three officers), firm-level SLAs, and systems that make the process excellent independent of the person. [§ 13 person-bound service, § 11.4 key-person risk, § 13 standard work]

Q5. Why do fee sheets match competitors to the dollar while service reviews range wildly — what differentiator is never priced?

A: Fees are regulated/filed (commoditized), so competition moved entirely to unpriced service quality. The differentiator — responsiveness, problem-solving — can't be priced line-item, so it's won via reputation and relationship. Compete where price is fixed: make service superiority visible and provable. [§ 2.3 regulated price competition, § 6.3, § 1.1]

Q6. Why do commercial deals carry 10× fees at 4× staff time — which book would we grow if honest?

A: Commercial — the fee-to-effort ratio is better. But commercial volume is thin and relationship-concentrated, while residential is flow business. Honest answer: grow commercial deliberately (dedicated officer, attorney relationships) while running residential as efficient volume. [§ 1.3 margin analysis, § 1.1 focus, § 2.1]

Q7. Why do agents complain about title fees to clients while choosing the cheapest option themselves — what does that game do to service investment?

A: It's blame-shifting theater: the agent gets low fees for their clients while performing advocacy. It compresses your margin to the point where service investment is irrational — which then justifies their complaints. Break the cycle by making service differences visible to agents directly (closing-experience metrics). [§ 13 game dynamics, § 11.3, § 6.3]

Q8. Why do error rates concentrate in refinance surges — volume should make us better, not worse?

A: Surge economics: volume arrives faster than capacity can scale (hiring/training lag), so experienced staff rush and overtime degrades attention. Volume improves quality only within capacity; beyond it, quality falls off a cliff. Surge playbook: overflow staffing agreements and error-rate-triggered intake throttling. [A4, § 4.2 capacity, § 6.3 SPC under overload]

Q9. Does same-day recording win business or become invisible once expected — which files cite it?

A: It wins business exactly once (when you first offer it), then becomes a qualifier — expected, invisible, unpriced. That's the order-qualifier treadmill: yesterday's winner is today's floor. Keep finding the next service edge; the treadmill never stops. [§ 1.1 OrderQualifier migration, § 13, § 6.1]

Q10. When wire fraud spikes, why does client protection depend on the officer's diligence rather than a system?

A: Because nobody built the system: verification callbacks, secure portals, and warning scripts exist as individual habits, not enforced process. Fraud resistance is exactly what process design is for — mandatory callback protocols with documented completion, portal-only wire instructions.

[§ 6.3 poka-yoke, § 14 security, § 13 standard work]

76. Real estate appraisal

Q1. Non-lender work carries better fees and steadier demand, yet every system is built for lender volume. What if effort followed margin?

A: You'd rebuild around private work: marketing to attorneys/CPAs/homeowners, consumer-friendly ordering, and report formats for lay readers. The lender infrastructure (AMC logins, UAD compliance) is a sunk-cost anchor. Run private work as a separate line with its own process. [§ 1.1 structural decisions, § 1.2 focus, § 1.3]

Q2. Revision requests cluster by AMC, not by appraiser. Where does quality control actually live?

A: In the AMC's review checklist, not your work: some AMCs run automated rule-checkers that bounce reports on technicalities regardless of valuation quality. Track revision cost per AMC and price/route accordingly — a high-revision AMC is a high-cost client regardless of fee. [§ 11.2 channel cost, § 13 Goodhart (checklist QC), § 1.3]

Q3. When volume drops, why do our best appraisers leave for assessment districts — and the book concentrates in whoever's left?

A: Because your best appraisers have the most options: government assessment jobs offer stability exactly when private volume is volatile. The book concentrating in the remainder is adverse retention. Counter-cyclical retention (retainers, salary floors) for your top tier costs less than rehiring at the next up-cycle. [§ 13, § 11.4 cycle resilience, § 4.4]

Q4. Does our turn-time reputation win the AMC panels we want or the volume we can't staff — which arrived?

A: Check which arrived: fast reputation typically wins volume-first AMCs (they sort by speed and fee), which flood you with commodity orders — the volume you can't staff. If you wanted quality panels (private banks, complex work), speed was the wrong signal to lead with. [§ 1.1 positioning signals, § 13, § 2.1]

Q5. Why do complex properties take 3× hours at 1.5× fees — and why do we keep accepting?

A: Because the fee schedule prices by form, not difficulty, and you accept from schedule-fear (an empty week feels worse than a bad fee). Complexity pricing is defensible — clients pay for competence on hard properties. Quote complexity premiums and let the schedule breathe. [§ 2.3 value pricing, § 1.3 opportunity cost, § 13]

Q6. Lenders demand speed; AMCs demand price; revisions come from the intersection. Who is the report for?

A: The underwriter's checklist, functionally — the report is written to pass automated review, not to inform a lending decision. You've optimized for the gatekeeper's gatekeeper. Where possible, cultivate direct-lender and private clients who read reports as opinions again. [§ 13 principal-agent chains, § 6.3, § 1.1]

Q7. If a court reviewed our last fifty reports, would it find opinions or defensibility — which did the fee pay for?

A: The fee paid for defensibility (form compliance, comp grids); a court tests opinion quality (reconciliation logic, judgment). Litigation work pays 2–3× precisely because it buys the opinion. The gap between your lender product and your litigation product is the gap between defensibility and thinking. [§ 6.1 quality dimensions, § 2.3, § 14 liability]

Q8. Why do comparable-selection disputes resolve in our favor 90% of the time, yet we keep spending the hours?

A: Because the dispute process has no cost to the requester: underwriters/Stip reviewers fire queries freely, and you pay the response labor. Track stip-hours per client; some AMCs' business model includes your free rebuttal labor. Price it (revision fees) or deprioritize the offenders. [§ 11.3 incentive misalignment, § 1.3, § 13]

Q9. When we train a trainee, why do they get licensed and leave at profitability?

A: Because licensure transfers their value to the open market while your comp stays at trainee logic. The training investment needs a retention structure: multi-year comp progression, signed training agreements (where enforceable), or a partnership track visible from day one. [§ 13 incentives, A9 learning curve economics, § 11.4]

Q10. Why does desktop/hybrid appraisal adoption lag client requests — whose liability discomfort sets the pace?

A: Yours: desktops shift inspection risk (you rely on third-party data collectors), and your E&O instinct resists. But clients request it because speed/cost win for low-risk properties. Adopt selectively (cookie-cutter properties first) with documented data-verification steps — the liability is manageable, the market shift isn't. [§ 11.4 risk, § 12 digital ops, § 2.1]

77. Convenience stores & gas stations

Q1. Why do inside margins concentrate in five categories while shelf space says otherwise?

A: Because shelf space was allocated by history and slotting deals, not margin-per-facing. Run space-to-margin analysis: the five categories (usually beverages, tobacco, snacks, lottery, energy) deserve the prime footage; the long tail should shrink or justify itself as traffic drivers. [§ 10 layout economics, § 5.2 ABC, § 1.3]

Q2. Commuters bought those gallons before the loyalty program existed. What did the program change besides price — what if we killed it?

A: Possibly nothing except margin donation: if members were already captive commuters, the program discounts existing demand. Kill-test on a subset (or pause enrollment): watch gallon retention. Many fuel programs only pay via the data/inside-attach, not gallons. [§ 2.3 loyalty economics, § 13, § 1.3]

Q3. Best managers run highest payroll % and best net profit. Which does the bonus plan measure?

A: If it measures payroll %, it punishes your best managers — they staff adequately (higher payroll) and reap sales, shrink control, and cleanliness that convert to profit. Bonus on net profit or controllable margin, never on a single cost ratio that confuses investment with waste. [§ 13 Goodhart, § 1.3 balanced metrics, § 4.4]

Q4. Why does foodservice succeed at two locations and fail at three identical ones — what variable have we refused to measure?

A: Execution and local demand: the same program dies under weak managers (freshness discipline, speed) or wrong demographics (no lunch traffic). Measure manager-level foodservice P&L and daypart traffic by store before rolling out further — the program isn't portable; the conditions are. [§ 6.3 stratification, § 10 location demand, § 13]

Q5. Are we a fuel business with a store, or a store with pumps — which survives EV adoption in our zip codes?

A: Check your EV adoption curve and inside-sales margin share: if inside profit already exceeds fuel margin, you're a store with pumps — and your capex (foodservice, seating, charging) should follow that identity. The zip-code EV trajectory sets your timeline, not your preference. [§ 1.1 identity/structural, § 2.2 leading indicators, § 10]

Q6. Why do overnight hours lose money at some stores and profit at others a mile apart — what do the profitable ones sell?

A: Different demand: profitable overnights serve shift workers, truckers, hospitality — with coffee, food, and essentials; losers serve nobody but still burn labor and lights. Overnight is a per-store decision: match hours to the actual overnight customer, not to chain uniformity. [§ 10 micro-location demand, § 1.3, § 2.1]

Q7. When a competitor drops fuel prices, we lose inside sales more than gallons. What are customers choosing on?

A: The trip itself: fuel price determines which station they stop at; the inside purchase follows the stop. You're not losing store customers — you're losing the traffic that feeds the store. Fuel price is your traffic acquisition cost; manage it as marketing, not margin. [§ 2.3 traffic economics, § 1.1 OrderWinner (location+price), § 13]

Q8. Fuel drives 70% of traffic and 30% of profit. Which number do site decisions use?

A: They should use both in sequence: fuel traffic is the funnel, inside conversion is the profit. Site decisions based on traffic counts alone ignore conversion quality (demographics, format); profit-based alone starve the funnel. Score sites on traffic × expected conversion. [§ 10 location models, § 1.3, § 2.1 funnel]

Q9. Why does shrink concentrate in products staff can consume — which policy have we never enforced?

A: The employee-consumption policy: drinks, snacks, and lottery are consumable/concealable, and "free shift drinks" norms blur into grazing. Enforce a written policy (paid items logged, designated free items) plus camera coverage of the cooler aisle — shrink follows ambiguity. [§ 14 internal controls, § 13, § 5.1 shrink as inventory variance]

Q10. When we remodel, why does lottery revenue jump more than grocery — which customer did the remodel attract?

A: The impulse/entertainment customer: brighter, cleaner stores lift immediate-gratification categories first; grocery basket growth requires trust (freshness, prices) that remodels don't confer. If the remodel was supposed to build grocery share, the follow-through is assortment and pricing, not lighting. [§ 6.1 tangibles, § 2.1, § 10 layout]

78. Clothing boutiques

Q1. Why do trunk shows sell volume to existing customers instead of acquiring new ones, event after event?

A: Because your invite list is your customer file: trunk shows are retention events dressed as acquisition. That's fine if priced as retention (they do lift frequency) — but stop judging them on new-customer counts, and build separate true-acquisition events (partner cross-promotions, neighborhood pop-ins). [§ 2.3 event economics, § 13, § 2.1]

Q2. Returns concentrate in fast, correct online shipments. What is the customer actually returning?

A: The photo-to-mirror gap: fit uncertainty and expectation mismatch, not logistics failure. Reduce with better fit information (model measurements, true-size guidance, video) and consider the bracket-buying reality — some returns are the cost of online apparel; price them into margin. [§ 6.3 expectation gap, § 11.1 SCOR Return, § 1.3]

Q3. Instagram items sell out online while fitting-room sellers never photograph. Which inventory does the buyer understand?

A: The online algorithm, not the local woman: buying follows scroll-appeal (photogenic, trend-forward) while your floor customer buys fit and flattery. Split the buy: online-exclusive photogenic capsule vs. floor-validated core — or reconcile by photographing the fitting-room winners properly. [§ 2.1 channel demand split, § 6.1, § 13]

Q4. When a carried brand goes DTC, why do we keep it on the floor at the same terms?

A: Inertia and fear of assortment gaps — but the brand now competes with you using your floor as its showroom. Renegotiate (exclusivity, better margin, MAP enforcement) or replace: a DTC brand at retail terms is a landlord relationship where you pay the rent. [§ 11.2 supplier power shift, § 11.3 channel conflict, § 2.3]

Q5. Would a first-time customer call us curators with a point of view, or a rack of brands available anywhere?

A: Test it: if customers can name what's distinctly you (a silhouette, a mix, a price-point thesis), you're curators; if the floor reads as "brands," you compete with the internet on its terms. Curation is the only defensible boutique position — edit harder. [§ 1.1 positioning/focus, § 6.1, § 13]

Q6. Why do part-timers outsell full-timers per hour — what does the full-time role contain?

A: Non-selling labor: receiving, merchandising, scheduling admin fills full-timers' hours, diluting their per-hour selling. Per-hour comparisons are unfair across different job contents — compare per-hour-on-floor, and consider whether full-time roles are carrying operations that deserve their own line. [§ 4.4 work content analysis, § 13 Goodhart, § 1.3]

Q7. Does personal styling build wardrobe loyalty or discount the full-price floor — what does stylist-client repeat data show?

A: The repeat data decides: if styling clients return at full price with higher basket sizes, styling is a loyalty engine; if they return only for styled-sale events, you've trained them to wait for curation-with-discount. Structure styling as a service (fee or minimum) rather than a discount channel. [§ 2.3 service design, § 1.3 BalancedScorecard customer perspective, § 13]

Q8. Full-price comes from regulars; the sale rack clears to strangers. Which does the buying calendar serve — what if we stopped marking down?

A: It serves regulars if the buy is disciplined; the sale rack exists because the buy overshot (strangers harvest your mistakes). Stopping markdowns without fixing the buy means dead stock. Fix the buy depth first (smaller initial buys, reorder winners), then markdowns shrink naturally. [§ 5.1 buy depth, § 5.2 newsvendor logic, § 13]

Q9. Why does December fund the year while January decisions assume every month is December?

A: Recency bias in budgeting: annual buys and rent commitments get set in the glow of peak season. Build the cash-flow calendar first (12 monthly columns, actual seasonality), then size commitments to the troughs, not the peak. December is 30% of revenue and 0% of predictability for March. [§ 2.1 seasonality, § 4.1 working capital, § 13 recency]

Q10. Markdowns clear the sizes we overbought, not the styles that failed. What is buying protecting?

A: The buyer's style conviction: size-curve errors are operational (fixable with size-curve data); style failures are judgment errors nobody wants to face, so failed styles linger full-price hoping. Review sell-through by style explicitly — the data absolves or convicts the buy, depersonalized. [§ 13 ego in forecasting, § 5.1, § 2.2 sell-through data]

79. Liquor stores

Q1. Why does shrink concentrate in minis at the register — the highest-visibility spot — and what's never changed?

A: The layout: minis are pocket-sized, high-value, and positioned where staff attention is lowest (register blind spots) — visibility to customers isn't visibility to staff. Move minis behind the counter or into locked/attended displays; the shrink map is a layout critique. [§ 10 layout, § 14 loss prevention, § 5.1]

Q2. When allocated products arrive, why do they go to whoever's in the store that day — what's that policy costing?

A: Your best customers' loyalty: allocated bottles are relationship currency, and random allocation spends it on strangers. A waitlist/priority system (top customers get first call) converts scarcity into retention. Random allocation is a fairness story that costs you your best accounts. [§ 2.3 allocation as loyalty tool, § 13, § 8.2]

Q3. Highest-margin craft items gather dust while low-margin domestic cases fund the store. Which does the shelf serve?

A: The shelf serves velocity habit: domestics get the cold box and the front; craft gets the warm back wall. Craft margin needs merchandising (staff picks, tasting notes, end-caps) — margin without visibility is theoretical. Rebalance prime space by margin-per-facing × realistic velocity. [§ 10 layout, § 1.3, § 2.3]

Q4. When Total Wine opens nearby, we lose bottom-shelf volume but keep the allocated-bourbon crowd. Which business are we in?

A: The relationship/allocation business: commodity volume is gone permanently (they win price and selection), but expertise, allocation access, and convenience remain yours. Pivot the store: deepen the allocated/craft identity, accept the volume loss, and stop price-fighting a battle that's over. [§ 1.1 positioning, § 11.4 competitive shock, § 2.1]

Q5. Customers ask for recommendations, then buy the shelf-talker brand anyway. What is expertise worth at the shelf?

A: Nothing at the shelf — the talker (paid marketing) beats your verbal advice at the decision moment. Move your expertise into the physical decision point: staff-pick tags with names and reasons ("Marcus's value pick") compete with talkers; your voice at the register is too late. [§ 13 point-of-decision, § 6.1, § 2.3]

Q6. Why do delivery orders skew premium while walk-ins buy value — which inventory does the buyer stock for?

A: Check the stockout data: if premium SKUs run thin because the buy follows walk-in velocity, you're underserving the channel with better basket economics. Stock for both channels explicitly — delivery customers are a different, richer demand pool using the same shelf. [§ 2.1 channel segmentation, § 5.1, § 1.3]

Q7. Why do wine-club members churn at month 6, right after the novelty bottles run out?

A: Because the club's promise was discovery, and your curation runs out of surprises — months 1–5 are highlights, month 6 is filler. Fix the curation engine (deeper supplier relationships, themed arcs, member-preference tracking) or restructure as quarterly shipments to slow the novelty burn. [§ 2.3 subscription design, § 13 novelty decay, § 11.2]

Q8. Does the tasting calendar build the base or entertain the same 30 regulars — what does attendee purchase data show?

A: The purchase data decides: if attendees' 30-day spend lifts and occasional new faces convert to regulars, tastings work; if the same 30 drink free monthly with flat spend, it's a social club you're sponsoring. Cap frequency for repeat attendees or charge a tasting fee redeemable on purchase. [§ 2.3 event ROI, § 13, § 8.2 fences]

Q9. What does price-matching confess: convenience retailer or specialty merchant — which do allocated-bottle customers believe?

A: Matching confesses convenience retailer (competing on price), while your allocated customers believe specialty merchant (access and expertise). The confession undermines the belief — merchants don't match; they justify. Drop the matching, invest in the access narrative. [§ 1.1 positioning consistency, § 6.1, § 13]

Q10. Why do Sunday hours generate best per-hour sales and most staffing complaints?

A: Because Sunday demand is concentrated and high-intent (pre-dinner, event runs) while staffing it requires weekend premiums and breaks your full-timers' rotations. The economics usually justify Sunday — solve the staffing with Sunday-specific part-timers and premium pay, priced against that per-hour revenue. [§ 7 labor scheduling, § 2.1 daypart, § 4.4]

80. Hardware stores

Q1. Why do rental and key-cutting — best margins — occupy the worst floor space?

A: Because floor space was allocated by product departments, not service margins: rental counters got the back wall when the store opened and never moved. Margin-per-square-foot analysis says move services forward — they're also your traffic-drivers and expertise showcases. [§ 10 layout economics, § 1.3, § 13 legacy decisions]

Q2. Which customer returns: the project planner or the problem-solver — which does staffing serve?

A: The problem-solver returns (leak, broken part, urgent need) — and they need floor expertise immediately. Staffing usually serves the planner (project desk) while problem-solvers wander aisles. Staff the aisles at peak problem hours (weekend mornings) with your best diagnosticians. [§ 2.1 demand type, § 6.3 responsiveness, § 7 scheduling]

Q3. Why does commercial delivery lose money per stop while being the stated reason contractors stay — has it been repriced?

A: It's an unpriced retention cost: delivery buys contractor loyalty (they stay for the convenience) but runs negative per-stop. Reprice honestly: free delivery over order minimums that make stops profitable, or a monthly delivery subscription. The loyalty value doesn't require the loss. [§ 2.3 pricing structure, § 1.3, § 11.2]

Q4. Why does our advice advantage show in surveys but not basket size — what would make expertise billable?

A: Because advice is free and unbundled: customers extract diagnosis, then buy the cheapest channel. Billable paths: project consultation fees (credited to purchase), installed sales (we do it), or membership (pro accounts with priority service). Expertise unpriced is expertise donated. [§ 2.3 service pricing, § 6.1, § 13 showrooming]

Q5. When we match online prices, why do customers still showroom us — which categories have we silently conceded?

A: Because matching removes the price objection but not the habit: showrooming persists where delivery convenience wins (heavy, bulk) or assortment wins (long tail). Concede explicitly in those categories (minimize inventory) and dominate where immediacy rules (repair parts, project-completion items). [§ 1.1 focus, § 2.1 channel behavior, § 5.1]

Q6. Does the special-order desk build loyalty or expose inventory gaps — what does special-order-to-regular conversion show?

A: Track it: customers whose first special order succeeds (right part, fast, communicated) become loyal — you solved what nobody stocks; failures (long silence, wrong part) confirm the big-box alternative. The desk is a loyalty machine only with proactive status communication. [§ 8.2 blueprint, § 2.3, § 6.3 responsiveness]

Q7. Why do our most knowledgeable employees gravitate to the service counter — the lowest-revenue square footage?

A: Because that's where the interesting problems are — experts seek diagnostic work, not aisle-facing. The fix isn't moving them; it's monetizing their position: the counter drives attach sales (parts for the diagnosed fix). Track counter-attached revenue before calling it low-revenue space. [§ 13 motivation, § 1.3 attribution, § 10]

Q8. Contractors open house accounts, use them a year, then pay cash at the big box. What did the account offer?

A: Apparently only credit — which they didn't value. Accounts that retain offer job-lot pricing, delivery priority, dedicated reps, and monthly consolidated billing (their bookkeeper cares). If your account is just a charge card, the big box's 5%-off card beats you. [§ 2.3 B2B value design, § 11.2, § 13]

Q9. Seasonal buys sell through some years and sit in others under the same forecast. What variable is missing?

A: Weather timing: seasonal sell-through depends on when spring arrives, first frost, snow events — not just average demand. Layer weather-based triggers into buy commitments (smaller initial, weather-triggered reorders) or negotiate supplier return terms on seasonal stock. [§ 2.1 weather-driven demand, § 5.2 newsvendor, § 11.2 terms]

Q10. When a big box opens two miles away, why do we lose project sales but keep emergency sales — which paid the rent?

A: Project sales (planned, price-shopped, bulk) migrate to the box; emergency sales (urgent, nearby, expertise-needed) stay. If projects paid the rent, the rent model must change: double down on the emergency/service identity (plumbing/electrical repair depth, contractor speed) where proximity wins. [§ 1.1 OrderWinner proximity/urgency, § 11.4 competitive shock, § 10]

81. Furniture stores

Q1. If we charged online-brand shoppers for the showroom service, what would it cost us to keep giving it away?

A: Compute: showroom hours $\times$ staff cost + floor-space cost $\div$ conversion rate of shoppers who buy from you. Showrooming losses are real but charging admission kills traffic; better: capture identity at entry (design consult signup) and price-match-with-service bundles that monetize the floor. [§ 13 showrooming, § 2.3, § 10 space economics]

Q2. When manufacturers raise prices, why do we hold ours until floor tags embarrass us — what does the lag cost?

A: The lag costs the spread $\times$ volume sold at stale prices — typically a full margin point for a quarter. The cause is tag-changing labor and fear of being first to move. Build a cost-increase pass-through policy (30-day lag max, automatic retag cycle); furniture customers accept cost-driven increases when stated plainly. [§ 11.3 price fluctuation, § 0.3 governedBy, § 2.3]

Q3. Delivery damage clusters on highest-end pieces — handling, expectation, or the subcontractor's routing?

A: Decompose claims by carrier and piece type: premium pieces often ride the same sub-contracted trucks as everything else (no white-glove protocol), and premium customers document damage more. Route high-value deliveries through a defined white-glove process with inspection sign-off at threshold. [§ 11.2 carrier management, § 6.2 external failure, § 8.2]

Q4. Why do mattress sales carry the store's margin while salespeople avoid that department?

A: Because mattresses feel like a different sale (higher pressure, sleep-trial complexity, comparison-shopped customers) and comp plans often underweight them vs. case goods. Fix the comp and the training: the margin lives where your staff won't go. [§ 13 incentives, § 4.4 skill, § 1.3]

Q5. Does design consultation sell furniture or give away the expertise — what do consultant-attached tickets show?

A: Compare ticket sizes: consultant-attached sales typically run 2–4 $\times$ walk-in tickets with lower discounting. If yours don't, consults are free advice sessions. Structure: consultation fee credited against purchase — filters serious buyers and prices the expertise. [§ 2.3 service pricing, § 1.3, § 13]

Q6. Why do financing approvals close sales that default at rates our margin can't absorb — who owns that trade-off?

A: Nobody — sales gets credit for the close; the finance company or your receivables eat the default. Assign ownership: track default rate by salesperson and promotion; if a promo (no-credit-needed events) manufactures defaults, kill it. The sale isn't complete at signature. [§ 13 incentive misalignment, § 11.4 credit risk, § 1.3]

Q7. Our best sales month coincides with worst gross margin, every year. Who decided that was acceptable?

A: The calendar did: your biggest month is a clearance/holiday event (volume bought with discounts). It's acceptable if it clears aged inventory at recovered cost; corrosive if it trains customers to wait. Separate the analysis: event margin on aged stock (good) vs. event discounts on current stock (bad). [§ 8.2 event yield, § 5.1 aging, § 13]

Q8. Customers visit three times before buying. What happens on visit two that sales doesn't touch?

A: The household negotiation: visit two is where they bring the spouse or the measurements — the decision shifts from taste to logistics and budget. Your process treats visit two as a fresh browse. Capture visit-one details, prep visit-two materials (room-fit checks, fabric samples), and recognize the returning couple as mid-funnel, not new traffic. [§ 8.2 blueprint, § 13 decision journey, § 2.3]

Q9. When we discount a floor model, why does its regular-price version stop selling?

A: Anchoring: the tagged floor model becomes the reference price for that item in customers' minds — the new one at full price reads as overpriced relative to "the same sofa" on clearance. Isolate floor-model sales (different SKUs visibly, "as-is" framing) so the anchor doesn't transfer. [§ 13 anchoring, § 2.3, § 10 floor management]

Q10. Why do special orders (8–12 week lead) give best margins and worst reviews — which do customers remember?

A: They remember the wait — especially the unmanaged wait (silence weeks 3–10). Margin is earned at order; reviews are lost in the silence. Fix the wait experience: proactive status updates at set intervals convert the same delay from neglect into process. [§ 8.2 blueprint, § 6.3 responsiveness, § 13 peak-end]

82. Independent pharmacies

Q1. Prescription volume grows while margin shrinks. Which does the business plan assume continues?

A: If it assumes volume growth saves it, the plan is broken: PBM economics compress per-script margin structurally — volume growth on shrinking unit margin is running faster downhill. The plan must pivot margin to clinical services, compounding, and front-end — or admit the script business is a traffic engine only. [§ 1.1 structural reality, § 1.3, § 2.1]

Q2. Why do immunization clinics generate goodwill that never converts to front-end sales?

A: Because the visit frame is clinical: patients enter, get the shot, leave — no shopping journey exists. Create one: post-shot 15-minute observation with relevant OTC recommendations (immune support, the season's needs) and front-end vouchers. Goodwill without a path stays goodwill. [§ 8.2 blueprint, § 2.3 conversion design, § 13]

Q3. What do staffing hours fund: a healthcare provider or a retailer renting a license — which do patients experience?

A: Count the hours: if pharmacist time goes to counting and verification while technicians could handle flow, patients experience a retailer. Provider status requires consultation time — staff for counseling availability (tech-driven workflow) so the pharmacist's face time is the product. [§ 4.4 labor allocation, § 1.1 positioning, § 6.3]

Q4. Why does DIR-fee exposure surprise us every quarter when the contracts never change?

A: Because DIR fees are retrospective (calculated months later on performance metrics), and nobody models the accrual: the surprise is a reporting design failure, not a contract surprise. Build a DIR reserve model per PBM contract using your performance data — the fee becomes a forecast line, not a quarterly ambush. [§ 2.2 forecasting, § 4.1 accruals, § 13 denial]

Q5. Why do compounding and specialty — our best margins — depend on two prescriber relationships?

A: Because specialty scripts flow from prescribers, not patients, and you never institutionalized the channel. Two relationships = two failure points carrying your margin. Systematize: prescriber-liaison visits, outcome reporting to their offices, and widen to 10+ prescribers before one retires. [§ 11.4 concentration risk, § 11.2 channel development, § 4.3 series reliability]

Q6. Does delivery retain patients or absorb cost — what does the retention differential show?

A: The differential answers it: delivery users typically show meaningfully higher retention (convenience lock-in, especially elderly/chronic patients). If confirmed, delivery is a retention product — price the route economics against retained LTV, not per-trip cost. [§ 2.3 retention economics, § 1.3 BalancedScorecard customer perspective, § 10 routing]

Q7. When a prescriber's e-prescribing default switches to a chain, why do we learn from declining volume, not the relationship?

A: Because you have no early-warning channel: the default switch happens in the prescriber's software, invisible to you until scripts stop arriving. The relationship fix: regular prescriber-office contact (your techs and their staff talk weekly anyway — formalize it) so default changes surface immediately. [§ 2.2 leading indicators, § 13 relationship maintenance, § 11.2]

Q8. Why do front-end sales decline as clinical services grow — one identity chosen, or both lost?

A: Clinical growth consumes the pharmacist and the floor's attention, and the front-end withers unmerchandised. It's not identity conflict — it's capacity conflict. Assign front-end ownership (a lead tech with margin KPIs) so clinical growth doesn't quietly starve your highest-margin square footage. [A3 resource contention, § 10 layout, § 1.3]

Q9. When a PBM cuts reimbursement below acquisition cost, why do we keep filling — what's the actual calculation?

A: The calculation is fear: losing the patient's other scripts and the relationship. But negative-margin scripts compound; the real math compares script-level loss vs. whole-patient LTV. For chronic below-cost fills, document and escalate (state protections, MAC appeals) or transfer the patient gracefully. [§ 1.3 unit economics, § 11.3 PBM power, § 13]

Q10. Why do adherence-packaging patients stay for years while walk-ins defect to mail order silently?

A: Packaging creates a service bond (weekly dependency, pharmacist relationship, complexity only you manage); walk-in scripts are commodities where mail order's convenience wins. The lesson: convert commodity patients into service relationships (sync programs, packaging, med reviews) before mail order's convenience claim lands. [§ 2.3 switching-cost design, § 6.3, § 8.2]

83. Florists

Q1. Why do wire-service orders fill the shop at margins that can't fund a designer — and why keep accepting them?

A: Because wire orders fill idle capacity (any contribution beats empty hours) — but they set a ceiling: a shop full of wire work has no capacity for profitable direct orders. Set a wire-order cap (accept only below X% of capacity) and invest the freed capacity in direct channels. [A2 opportunity cost, § 8.2 capacity allocation, § 1.3]

Q2. Why do we have sold-out days AND waste in the same Valentine's week, every year?

A: Because demand spikes unevenly across the week and across SKUs: you stock out of roses on the 13th while mixed arrangements die on the 15th. Newsvendor discipline per SKU-day: pre-book commitments drive the buy, and walk-in stock is sized by day-of-week-within-peak history. [§ 5.2 Newsvendor, § 2.1 spike decomposition, A5]

Q3. Which failure mode shows in the books: artists with a retail problem, or retailers with a perishability problem?

A: Check waste % and labor %: high waste = perishability mismanaged (retail failure); low waste but high labor-per-arrangement = artistry unfunded (pricing failure). Most florists are artists with a retail problem — the fix is recipe costing and engineered arrangements, not more art. [§ 1.3, § 5.1 perishables, § 6.1]

Q4. When stem costs spike 30% at Valentine's, why does margin get crushed exactly at peak volume — which lever is unpulled?

A: The price lever: fear of holiday sticker-shock keeps prices flat while costs spike — but Valentine's buyers are the least price-sensitive of the year (deadline, obligation, no comparison shopping). Raise holiday prices 10–15%, adjust recipe sizes, or pre-buy contract stems. You've never tested the ceiling on your most inelastic demand. [§ 2.3 inelasticity, § 11.2 pre-buy, § 13]

Q5. Why do delivery zones include money-losing addresses drawn by habit?

A: Zones were drawn when the shop opened and never repriced against actual driver time. Recompute per-zone delivery cost (time × wage + vehicle); far addresses need zone surcharges or minimums. Habit zones quietly tax every profitable delivery. [§ 10 routing/zone economics, § 1.3, § 13 legacy policy]

Q6. Why do grocery bouquets take casual volume while custom work holds — which business did we think we were in?

A: You thought "flower business"; the market split it: commodity sentiment (grocery wins on price/convenience) vs. designed emotion (custom holds). The middle is gone. Commit to the custom/event identity and let the grocery have the casual stem — or build a grab-and-go cooler that competes on convenience. [§ 1.1 positioning, § 2.1 market bifurcation, § 13]

Q7. Does Instagram drive orders or admiration — what does order geo vs follower geo show?

A: Usually admiration: flower followers are global, orders are 10-mile. If follower-geo ≫ order-geo, Instagram is a portfolio, not a channel. Local conversion lives in Google (maps, reviews), funeral-home and venue relationships, and neighborhood visibility. Measure channel by delivered-order attribution. [§ 2.2 channel attribution, § 13 vanity metrics, § 10]

Q8. Why do wedding consults book a year out while everyday delivery erodes — which does the studio layout serve?

A: The layout serves weddings (consultation table, portfolio walls) while everyday orders get the phone-and-cooler treatment. Everyday delivery is your cash-flow base; it deserves its own visible operation (online ordering ease, same-day cutoffs, driver reliability). Don't let the glamorous line starve the annuity. [§ 1.1 focus balance, § 8.2, § 1.3]

Q9. Sympathy and wedding orders depend on funeral-home and venue relationships never formalized. What would formalizing be worth?

A: The value of your two highest-AOV channels' stability: formalize with preferred-florist agreements (reciprocal referral terms, package pricing, standing contact). Informal relationships die when the funeral director's niece opens a shop. Formalization converts luck into channel equity. [§ 11.2 channel formalization, § 11.4 relationship risk, § 2.3]

Q10. Subscription accounts pay late and churn fast, yet are our only predictable revenue. What would event work do to cash-flow risk?

A: Event work trades payment predictability for deposit-based cash (50% upfront) — better cash timing, worse volume certainty. A mix (events + subscriptions + standing corporate accounts) diversifies both risks. Current state: your "predictable" revenue is slow-paying and churning — predictable doesn't mean healthy. [§ 4.1 working capital structure, § 11.4 risk mix, § 1.3]

84. Jewelry stores

Q1. Why does the safe hold more inventory value than annual revenue — what turn rate was never set?

A: Turns below 1.0 mean you're running a museum with a cash-flow problem: aged inventory's opportunity cost exceeds most margin. Set turn targets by category (bridal 1.5, fashion 2.0, estate 0.8), age-band the stock, and melt/markdown the zombies. [§ 5.2 ABC/turns, § 4.1 working capital, § 5.1 aging]

Q2. Engagement buyers become best lifetime customers while marketing chases gift occasions. Which does the case layout reflect?

A: Probably the gift occasions (fashion cases front, bridal tucked back) — backwards. The engagement sale is your annuity entry (wedding bands, anniversaries, upgrades for decades). Put bridal at the experiential center and build the lifecycle capture (band, anniversary, push present). [§ 1.3 BalancedScorecard customer perspective, § 10 layout, § 2.3]

Q3. Why does repair lose money per ticket while generating the traffic that sells everything else — has repair been priced as marketing?

A: No — it's priced as a service and loses, while its traffic value goes uncounted. Add it up: repair visits × conversion-to-purchase × margin. If repair-as-marketing clears its losses (usually yes), formalize it; also raise repair prices toward cost-recovery — customers pay for bench trust. [§ 1.3 attributed value, § 2.3 loss-leader logic, § 6.3]

Q4. What survives the customer's next anniversary: the object we sold, or the occasion we attached it to?

A: The occasion — the object is metal; the memory is the product. That means the sale isn't complete at purchase: anniversary reminders, care services, and upgrade conversations at each milestone keep you attached to the occasion cycle. [§ 2.3 lifecycle marketing, § 6.1, § 13]

Q5. Why do custom clients refer constantly while case-sale clients never do — which does commission reward?

A: Commission usually rewards the immediate gross (case sale) while custom work (longer, consultative) creates the referring evangelist. Custom is co-creation — clients tell that story forever. Rebalance commission toward custom origination and track referral-attributed revenue by sale type. [§ 13 incentives, § 2.3 referral mechanics, § 1.3]

Q6. Does lab-grown vs natural mix reflect margin, demand, or the owner's conviction — which would the floor say?

A: The floor would say conviction or confusion: the mix decision is ideological in most stores, not data-driven. Run it by the numbers: lab-grown grows share and ticket-accessibility; natural holds value-story and older buyers. Stock both with segment-matched presentation and let sell-through arbitrate. [§ 2.1 demand data, § 5.1, § 13 ideology vs data]

Q7. Why do highest-value sales happen in December while highest-value relationships form in June — which does staffing serve?

A: Staffing serves December (seasonal help, extended hours) — correctly for revenue — but June's relaxed pace is when consultations deepen and custom work originates. Protect June's consultative capacity; the relationships formed in summer are the December tickets. [§ 2.1 seasonality of relationship vs transaction, § 7, § 13]

Q8. When gold spikes, estate-buying surges at margins new inventory can't match — which business do we staff for?

A: Usually neither deliberately: estate buying is counter-cyclical margin (spikes when retail slows) — a natural hedge you staff reactively. Staff and market the estate-buying line during gold spikes explicitly (appraisal events, advertising) — it's your best margin exactly when retail is hardest. [§ 2.1 counter-cyclical, § 2.3, § 11.4 hedge]

Q9. Why do appraisal services book solid at a decade-old fee schedule?

A: Because the fee was set when appraisals were a favor, and nobody repriced as insurance requirements made them mandatory. Booked-solid at stale prices = money left on the table with zero demand risk. Reprice to market (per-item rates, update-cycle pricing) — demand is inelastic here. [§ 2.3 inelastic pricing, § 1.3, § 13 legacy pricing]

Q10. When customers buy online and service in-store, why do we absorb warranty work on margin we never earned?

A: Because refusing feels like brand damage — but manufacturer warranty terms often don't obligate you for third-party purchases. Set the policy: service welcome at posted rates; warranty coverage only for your sales. State it cheerfully; the free-warranty habit is a leak, not a courtesy. [§ 0.3 governedBy policy, § 2.3, § 13]

85. Pet supply stores

Q1. Which identity does our welfare reputation depend on: pet store or pet-food store — what if we exit live animals?

A: Audit the reputation's source: for most independents, live-animal sales are small revenue but big identity ("a real pet store"). Exit risks losing the identity anchor; staying carries welfare scrutiny. Many thrive post-exit by substituting adoption partnerships — you keep the heart, lose the inventory risk. [§ 1.1 identity decision, § 14 triple bottom line (people), § 1.3]

Q2. Why do customers research premium food with us then subscribe to online autoship — what would retention pricing look like?

A: They harvest your expertise and buy on convenience economics. Retention pricing: your own autoship (same discount, plus free local delivery + the advice layer). You can't beat their logistics, but you can match the mechanism and add what they lack — the person who knows their dog. [§ 2.3 subscription counter, § 8.2, § 6.1]

Q3. Why do loyal dog-food customers have cats they feed from online — which aisle did we concede?

A: The cat aisle: your assortment, expertise display, and staff energy skew canine, so cat owners concluded you're a dog store. The customer proved loyalty transfers across pets — your merchandising didn't. Fix cat category visibility and staff knowledge; the household wallet is already won half. [§ 2.1 share-of-wallet gap, § 10 layout, § 13]

Q4. Why do raw/fresh customers — highest spenders — require freezer capacity we treat as an afterthought?

A: Because the freezer is infrastructure, not a department: it gets capex approval once and no management since. Highest-spend customers deserve the reverse: freezer capacity as a managed category (SKU turns, stockout tracking, expansion triggers). Stockouts here lose your best customers to the subscription services. [§ 5.1 cold-chain as constraint, § 1.3, § 4.2 capacity]

Q5. Adoption weekends fill the store with future customers who buy their first bag elsewhere. What are we absent from in the 72 hours after?

A: The panic-shopping window: new adopters buy everything within 72 hours (usually at a big box on the way home). Be in that window: adoption-kit bundles ready at the event, a "new pet" coupon book with 7-day expiry, and food-transition guidance attached to the adoption paperwork. [§ 2.3 timing/window capture, § 13, § 8.2]

Q6. Why does delivery serve customers within walking distance — what problem is it solving?

A: The heavy-bag problem: 30-lb dog food doesn't walk home pleasantly, and elderly customers can't lift it. Delivery within walking distance is actually your retention product for your oldest, most loyal segment. Price it as service (free over threshold) rather than apologizing for the geography. [§ 2.1 demand reality, § 6.3, § 2.3]

Q7. Grooming drives retail sales but retail never converts to grooming. Why does the funnel run one way?

A: Because grooming clients wait in-store (captive browsing = retail attach) while retail customers never encounter grooming (booked out, invisible, no counter offer). Reverse-flow fix: groomer introductions at retail checkout, "meet the groomer" events, first-groom discount with food purchase. [§ 8.2 cross-service blueprint, § 2.3, § 13]

Q8. When a brand launches its own autoship site, why do we keep stocking and recommending it while it harvests our customer list?

A: Because delisting feels like losing a seller — but you're funding your disintermediation. Options: negotiate (independent-channel pricing protection), de-emphasize (don't recommend; stock silently), or replace with a brand that protects the channel. Your recommendation is the asset they're monetizing. [§ 11.3 channel conflict, § 11.2 supplier terms, § 13]

Q9. Does nutrition expertise survive the phone price-check in the aisle — which categories does it still win?

A: It wins where the decision is complex and risk-laden: prescription-adjacent diets, allergy management, puppy/kitten stages. It loses on commodity adult maintenance food. Concentrate expertise merchandising where it wins; match price or bundle where it can't. [§ 6.1 expertise value by category, § 2.3, § 1.1]

Q10. When we price-match, why do baskets fall while traffic holds — which metric does the policy protect?

A: Traffic — the match keeps them coming but signals "we're the same as online," so the basket shrinks to the matched item (the reason-for-visit premium disappears into commodity framing). Protect basket: match on demand, but bundle (food + treats + toy pricing) so the trip stays a basket. [§ 13 framing, § 2.3, § 1.3]

86. Thrift & consignment**Q1. Why do donation goods carry better margins than consigned, yet we market to consignors?**

A: Because consignors bring curated volume (and become shoppers), while donations are erratic. But the margin says donations (100% COGS-free) fund you. Market both with the right frame: donations for mission/community marketing, consignment for quality curation. Know which feedstock actually pays. [§ 5.1 input economics, § 2.3, § 1.3]

Q2. Best categories (vintage, designer) depend on intake luck while staffing is built for processing volume. Which is the bet?

A: Processing volume, by default — but the margin is in the luck categories. Reduce the luck: cultivate dealer/estate relationships, consignor networks in affluent zip codes, and offer pick-up for quality estates. Intake quality is a sourcing strategy, not weather. [§ 11.2 sourcing structure, § 2.3, § 1.1]

Q3. Resellers find underpriced gems our staff missed. What is pricing missing?

A: Research time and brand knowledge: your pricers can't comp every item, so valuable pieces slip through at thrift prices. Fixes: triage protocol (unknown brands get looked up; vintage cues get flagged), or lean in — price for volume and let resellers be your velocity channel. [§ 6.3 inspection/knowledge limits, § 2.3, § 13]

Q4. Does payout structure (cash vs store credit) select for consignors or for inventory quality — which does it produce?

A: Cash attracts professionals with good inventory (they have options); credit-only attracts casual consignors (mixed quality) but retains spend in-store. Cash + credit-bonus hybrid gets quality AND retention: offer credit at 10–20% above cash value. [§ 13 incentive design, § 11.2 supplier selection, § 2.3]

Q5. Best consignors' items sell in days while 60% of consigned inventory never moves. Why haven't intake standards changed?

A: Because rejecting items is socially hard (the consignor is standing there) and standards are vibes. Write acceptance criteria (brands, condition, season) and train staff to decline gracefully — every accepted dud consumes rack space that sells-through would pay for. [§ 0.3 specifiedBy acceptance criteria, § 5.1 space opportunity cost, § 13]

Q6. Pricing staff vary 50% on identical items. What does the spread do to sell-through and consignor trust?

A: It degrades both: overpriced items stall (killing sell-through), underpriced ones vanish (consignors feel robbed when they notice). Build price guides by category/brand/condition with look-up tools, and audit pricing variance monthly — consistency is a process output. [§ 13 standard work, § 4.4 work measurement, § 6.3]

Q7. Why do online listings (eBay/Poshmark) out-margin the same items on the floor — which channel gets the good inventory?

A: Online wins on audience size (national collectors vs. local foot traffic). The good inventory should go online-first: triage at intake — collectible/brand items to online, commodity clothing to the floor. If the floor gets first pick by tradition, you're routing your best margin to your smallest audience. [§ 2.1 channel-value matching, § 10, § 1.3]

Q8. If the P&L testified: retail store or materials-recycling business with a storefront?

A: For many thrifts: recycling (rag-out, bulk sales, metal) carries the floor's losses. If so, optimize both honestly: the storefront as brand + mission + margin where it works, the recycling stream as the industrial base. The P&L's testimony should set the capex, not the self-image. [§ 1.1 emergent identity, § 1.3, § 14 circular economy]

Q9. Why do regulars visit weekly and buy nothing most weeks — what is the visit worth?

A: A lottery ticket: they hunt for the find, and the hunt IS the product. Weekly non-buying visits are engagement, not failure — they convert occasionally and bring others. Protect the hunt (fresh stock daily, treasure rotation); monetize with small-ticket impulse at checkout. [§ 13 engagement economics, § 6.1 experience, § 2.3]

Q10. Why do seasonal changeovers clear space and destroy sellable inventory simultaneously?

A: Because changeover policy is date-driven ("summer out July 1") not sell-through-driven: good items get clearanced or rag-out with the duds. Tier the changeover: seasonal turn moves only aged/low-turn stock; proven sellers cross seasons. Calendar convenience is costing you inventory. [§ 5.1 disposition policy, § 5.2, § 13]

87. Vape & smoke shops

Q1. Why do highest-margin products carry highest regulatory risk — which does the reorder cycle prioritize?

A: The margin, usually — reorders follow sell-through without a risk weight. Add a regulatory-risk factor to the buy: concentration limits on single-SKU/flavor exposure so a ban doesn't strand a third of your inventory. Margin without risk-adjustment is how vape shops die. [§ 11.4 regulatory risk, § 5.1 concentration, § 5.2]

Q2. Why does the wholesale side hustle out-earn a store location with zero management attention?

A: Because wholesale has better structural economics (no retail labor, volume B2B) and your attention follows the visible store. The side hustle is telling you where the scalable business is. Give it an owner, a quota, and a system — or sell it. Neglected profit centers decay. [§ 1.1 emergent strategy, § 1.3, D5 orphan activity]

Q3. Why do glass and accessories — our legal-safe margins — occupy the back of the store?

A: Because the front was built for the regulated volume drivers (vape disposables), and the safe margin inherited the back. As regulatory risk rises, the layout is backwards: merchandise toward the durable business (glass, accessories, lifestyle) while the front still monetizes the vape traffic. [§ 10 layout, § 11.4, § 1.3]

Q4. Does staff knowledge build the customer base or serve the existing one — which new segment has never walked in?

A: Likely the curious-newcomer segment (older, wellness-adjacent, CBD-curious) who find your store's vibe intimidating. Knowledge serving existing enthusiasts deepens loyalty; knowledge welcoming newcomers expands the base. Train the welcome script, not just the product lore. [§ 2.1 segment expansion, § 6.3 empathy, § 13]

Q5. Customers show brand loyalty to disposables that get banned. What happens to that customer when the SKU vanishes?

A: They churn to whoever has a substitute ready — brand loyalty transfers to availability. The defense: at ban-announcement, immediately curate the legal alternative (comparable flavor/nicotine profile) and contact your known buyers of that SKU directly. The transition moment decides if they stay yours. [§ 11.4 disruption response, § 2.3 retention, § 12 CRM]

Q6. Why do age-verification moments create our worst reviews — whose business are the one-star reviewers?

A: Underage or ID-less would-be buyers punishing your compliance — these one-star reviews are actually compliance documentation. Respond publicly and cheerfully ("we check every ID, every time") — the review you call worst is your license defense and your trust signal to legitimate customers. [§ 14 compliance as brand, § 6.1, § 13]

Q7. When card processors drop us, why does cash-only cost 20% of revenue but 0% of best customers?

A: Because best customers adapt (they know the drill, hit the ATM) while marginal/convenience customers vanish. The 20% you lose was your growth fringe, not your base. Mitigate: ATM on-site, cash-discount framing, and chase high-risk processor relationships before you need them. [§ 11.4 infrastructure risk, § 2.1 segment resilience, § 13]

Q8. Two locations stock identically while customer bases buy completely differently. What would per-store buy plans reveal?

A: That one store's bestsellers are the other's shelf-warmers: demographics differ (college town vs. suburb = disposables vs. glass). Per-store buys cut dead inventory and stockouts simultaneously. First quarter of store-level ABC analysis typically frees 15–20% of inventory value. [§ 5.2 ABC per location, § 2.1 local demand, § 5.1]

Q9. When a flavor ban hits, revenue dips exactly six weeks then recovers. What does the recovery consist of?

A: Substitution: customers cycle through resistance (stockpiling, complaints) then accept alternatives — the recovery is your remaining SKUs absorbing the demand. Knowing the six-week curve, you can manage the dip (pre-buy alternatives, customer communication) instead of panicking at week two. [§ 2.1 substitution dynamics, § 11.4, § 13]

Q10. Which inventory turns tell the truth: nicotine retailer diversifying, or accessories shop riding nicotine?

A: The turn $\times$ margin matrix tells it: if accessories turn slowly but carry margin while nicotine turns fast on thin margin, you're a nicotine retailer with an accessories hedge — today. Track the ratio's trend; when accessories' gross-profit share crosses nicotine's, your identity (and your regulatory exposure) has actually changed. [§ 1.3, § 2.2 trend analysis, § 1.1 identity]

88. E-commerce sellers

Q1. Our best-selling SKUs carry worst contribution margins after ads. Which number does the ad platform report?

A: ROAS on attributed revenue — not margin, not post-return reality, not cannibalized organic sales. Build contribution-per-SKU-after-ads (margin – ad cost – returns – fulfillment) and re-rank your catalog; platforms optimize for their revenue, not yours. [§ 13 Goodhart/platform metrics, § 1.3 contribution, § 11.3]

Q2. Why does the hero product fund a catalog of losers we keep restocking out of completeness instinct?

A: Completeness instinct is assortment vanity: the tail SKUs exist because a store "should" have them. Run contribution per SKU; cut or made-to-order the losers, and redeploy the working capital into hero inventory depth (which prevents the stockouts that actually cost you). [§ 5.2 ABC, § 4.1 working capital, § 1.1 focus]

Q3. Why do returns cluster in the size variant we photograph best?

A: Because the best photos create the most purchases of the variant with the worst fit accuracy — imagery outruns product reality. Either fix the variant's fit/sizing info (true measurements, fit notes) or accept it's your acquisition star with a return tax priced in. [§ 6.3 expectation gap, § 11.1 Return process, § 13]

Q4. A platform policy change could zero our traffic overnight. Which assets survive that morning?

A: Only what you own: the email/SMS list, brand search volume, supplier terms, and any off-platform revenue. Audit quarterly: what % of revenue could you regenerate from owned channels in 90 days? If under 20%, every marketing dollar should partly build owned assets. [§ 11.4 platform risk, § 4.1 intangible assets, § 1.1 structural]

Q5. When we raise prices, conversion holds but ad efficiency collapses — which do we act on?

A: Both are the same phenomenon: higher price filters out ad-driven cold traffic (price-sensitive clickers) while warm/organic converts fine. Act on contribution: if higher price $\times$ lower volume nets more profit, the ad efficiency loss is the price of margin. Judge by profit, not platform metrics.

[§ 2.3 price-volume trade-off, § 13 metric confusion, § 1.3]

Q6. When a supplier raises MOQs, why do we accept inventory risk rather than renegotiate — what's the switching cost?

A: Usually the real switching cost is lower than the inventory risk you're accepting — but renegotiation feels confrontational and switching feels slow, so absorption wins by default. Price the risk ($\text{MOQ} \times \text{failure probability} \times \text{markdown loss}$) and present it: suppliers flex MOQs for customers who quantify. [§ 11.2 supplier negotiation, § 5.1 risk, § 13]

Q7. Why do launches succeed with the email list and fail with paid traffic — which channel are we scaling?

A: Probably paid (it's dial-adjustable; the list grows slowly). But the list succeeds because it's warm trust; paid fails because cold traffic buys proven products, not launches. Launch to the list, harvest reviews, THEN scale paid on social proof. Sequence the channels by trust temperature. [§ 2.3 launch sequencing, § 13, § 2.1 channel trust]

Q8. Our best month (Q4) produces our worst cash flow. Which line assumed otherwise?

A: The inventory line: Q4 inventory is bought in August–October on cash, while revenue arrives in December and platform payouts lag further. The P&L celebrates; the cash flow starves. Model the cash conversion cycle by month — Q4 needs a war chest or a credit line sized to the buy, not the sales. [§ 4.1 working capital/cash cycle, § 2.1 seasonality, A5]

Q9. Does review velocity drive organic rank, or rank drive reviews — which did we just spend money on?

A: Both directions exist (a flywheel), but the spend determines which you bought: if you bought reviews directly, you bought velocity (risky, policy-violating); if you bought ads/placement, you bought rank. The durable flywheel is product quality $\rightarrow$ reviews $\rightarrow$ rank; money accelerates it but can't substitute. [§ 13 causation audit, § 11.3 platform mechanics, § 6.1]

Q10. Repeat rate looks healthy in aggregate while top-decile customers quietly churn. What is the aggregate hiding?

A: That your best customers are leaving while mediocre ones repeat: top-decile churn means your highest-LTV cohort hit a ceiling (outgrew your range, felt unvalued, got poached). Cohort by decile: if the top is churning, add VIP treatment and range expansion before the aggregate notices. [§ 2.2 cohort/decile analysis, § 13, § 1.3 BalancedScorecard customer perspective]

89. Trucking (owner-operators & small fleets)

Q1. Best drivers leave for 5-cent raises while retention budget targets sign-on bonuses. Who is the budget for?

A: Recruiting, not retention — sign-on bonuses are visible and attributable; retention raises feel like cost creep. But replacing a driver costs multiples of a 5-cent raise (recruiting, orientation, safety curve). Flip the budget: tenure-milestone raises beat sign-on bounties on pure math. [§ 13 incentive misallocation, § 1.3 turnover cost, § 4.4]

Q2. Which P&L line tells the truth: trucking company or finance company that owns trucks?

A: Compare margin from operations vs. the cost structure of equipment: if depreciation + interest dominate and operating margin is thin, you're running an asset-finance business where the truck note dictates decisions (take bad freight to cover the payment). Know it, because the fix differs (asset utilization vs. rate discipline). [§ 1.3, § 4.1 capital structure, § 1.1 identity]

Q3. Does factoring fund growth or hide a cash-conversion problem — what's the all-in cost per dollar?

A: Compute all-in factoring cost (fee + reserve holdbacks + admin) annualized — often 15–25% APR equivalent. If it funds growth into receivable-heavy customers, it compounds; if it papers over slow-pay customers you never renegotiated, it's a patch. Fix the terms or the customer mix, then quit the factor. [§ 4.1 working capital, § 11.2 terms, § 1.3]

Q4. Why does revenue per truck ceiling at 7 trucks no dispatch effort breaks — which function was never staffed?

A: The back office: at ~7 trucks, compliance, billing, maintenance coordination, and driver management exceed founder capacity — the founder becomes the constraint. The next hire isn't dispatch; it's operations/admin. Revenue per truck ceilings are org-chart problems wearing truck numbers. [A3 shifting bottleneck, § 1.1 structural, § 13]

Q5. We track detention precisely enough to resent it but never bill it. What would the annual invoice total be?

A: Tally it — detention hours × contractual rate typically sums to a shocking figure (often a truck's annual profit). Billing it costs some friction with shippers; not billing it costs the figure. Start with the worst offenders: documented detention + invoices + accessorial terms in the next rate agreement. [§ 0.3 measuredBy without billing, § 11.2 contract terms, § 1.3]

Q6. Why does maintenance cost per mile vary 40% across identical trucks — driver, route, or shop?

A: Stratify: driver behavior (idle time, harsh events — telematics shows it), route (mountains vs. flat, urban vs. highway), and PM compliance per unit. Usually driver + route dominate. Telematics-based driver coaching pays back fastest; the shop variance is rarest but check it last. [§ 6.3 stratification, § 12 telematics, § 4.3 maintenance]

Q7. When a shipper cuts volume, why do we learn from the load board, not the relationship — whose relationship was it?

A: The broker's or the rep's — not yours: if you never had direct contact with the shipper's transportation decision-maker, you were a capacity vendor, not a partner. Direct relationships with shippers (quarterly reviews, service data) turn volume changes into conversations instead of surprises. [§ 11.2 relationship depth, § 2.2 leading indicators, § 13]

Q8. Why do safety violations cluster in new drivers' first 90 days — what does orientation actually teach?

A: Paperwork, apparently — not your operation's realities: your routes, your customers' yards, your equipment quirks. Violations cluster because the learning happens on live freight. Structured finish-training (paired runs, route familiarization, 30/60/90 check-ins) compresses the violation window. [A9 learning curve, § 4.4 onboarding, § 14 safety leading indicators]

Q9. When spot rates spike, why does contracted freight block us from capturing it — which book do we regret?

A: You regret contracts in spot booms and spot exposure in busts — that's the portfolio trade. The answer is a deliberate mix (60–70% contract base, 30–40% spot flexibility) chosen in advance, not a posture that flip-flops with regret. [§ 8.2 yield/portfolio mix, § 11.4 market risk, § 2.1]

Q10. Highest-paying lanes produce worst net per mile after deadhead. Which number does dispatch optimize?

A: Gross rate per loaded mile — the visible number — instead of net per total mile including repositioning. Dispatch optimization needs the full-cycle metric $(\text{rate} \times \text{loaded miles} - \text{costs}) \div \text{total miles}$; the best-paying lane that strands your truck is a subsidy to the load board. [§ 13 Goodhart, § 1.3 net-mile economics, § 10]

90. Freight brokerages

Q1. Why does new-rep ramp take 14 months while training claims 6 — which does headcount planning use?

A: The 6, probably — which means you're perpetually under-staffed against reality. The 14-month truth is that book-building takes carrier relationships and shipper trust that training can't compress. Plan headcount on 14, and find the mid-ramp revenue bridges (house accounts, team structures) that survive it. [A9 learning curve reality, § 7 capacity planning, § 13]

Q2. Why do reps' books die when they leave despite the CRM holding every contact?

A: Because the CRM holds contacts, not relationships: the shipper trusted the rep's judgment and responsiveness, none of which lives in fields. Institutional defense: team coverage (two reps per account), documented shipper preferences/quirks, and company-level service reviews. [§ 13 relationship capital, § 11.4 key-person risk, § 12 CRM limits]

Q3. Why do we win RFQs at margins we'd reject on the spot board — which customer taught us which behavior?

A: Enterprise RFQ customers taught you to bid thin (volume promises, procurement pressure) while the spot board shows true market-clearing rates. RFQ margin discipline: set walk-away floors from spot data, and let procurement win someone else's capacity at loser's prices. [§ 2.3 bid discipline, § 11.3 procurement games, § 1.3]

Q4. Highest-margin loads come from smallest shippers while enterprise sets price and terms. Who is the pricing desk built for?

A: Enterprise (the visible volume) — while the margin lives in SMB shippers who need your expertise and have no procurement department. Rebalance the desk: enterprise gets efficiency treatment (automation, thin service), SMB gets margin treatment (service, expertise, loyalty pricing). [§ 2.1 segmentation, § 1.3, § 2.3]

Q5. When a carrier double-brokers our load, why does the cost land on our shipper relationship while the insurance claim lands nowhere?

A: Because double-brokering voids coverage chains (the actual carrier is unknown/unvetted), leaving you holding shipper trust and no recourse. Defense: carrier vetting rigor (authority age, inspection history), tracking verification, and contract teeth. One double-brokered disaster costs a year of margin. [§ 11.4 fraud risk, § 11.2 vetting, § 14 compliance]

Q6. If shippers described us: technology company with freight, or phone business with a website — which did we budget for?

A: The budget confesses: if 90% of spend is headcount and the TMS is a glorified spreadsheet, you're a phone business — fine, if your service justifies it. The risk is digital-native brokers commoditizing your service layer. Match budget to the identity you can actually win with. [§ 12 digital layer, § 1.1 positioning, § 1.3]

Q7. Why do carrier-relationship scores predict on-time performance better than vetting scores — which does coverage rely on?

A: Vetting scores check legitimacy (authority, insurance); relationship scores measure behavior (they answer your calls, they've hauled for you). Behavior predicts performance. Build the relationship layer systematically (carrier tiers, repeat-carrier preferences) instead of relying on marketplace roulette. [§ 11.2 carrier management, § 13, § 6.3]

Q8. When capacity tightens, why do we serve worst-paying shippers first out of habit — who re-ranks the board?

A: Nobody — coverage follows habit and relationship noise, not margin. In tight markets, your trucks (carriers) are the scarce asset: allocate to margin $\times$ relationship-value deliberately. A daily margin-ranked board in tight weeks is worth points of gross margin. [§ 8.2 yield allocation, A3, § 13]

Q9. Why do top-10 shippers hold 60% of volume and 30% of margin — which does the growth plan grow?

A: If the plan grows volume, it grows the concentrated, thin-margin enterprise book — doubling your dependence on your least profitable demand. Growth targeting mid-market shippers (where margin lives) diversifies both concentration and margin. Point the plan at the margin. [§ 11.4 concentration, § 1.3, § 1.1 deliberate growth]

Q10. Does tracking technology satisfy shippers or generate exception data nobody reads — which alerts changed an outcome?

A: Audit the alerts: if exceptions (late risk, temperature, route deviation) never trigger an action, the tracking is theater for the shipper portal. The value is in intervention: assign exception-handling ownership, measure saves. Tracking without response is expensive reassurance. [§ 12 sensing vs response, § 13 Goodhart, § 7 real-time control]

91. Moving companies

Q1. Last-minute bookings carry best margins and worst crew availability. Which does the calendar protect — what would a rush premium change?

A: The calendar protects early bookings at standard rates while rush demand pays the same and strains crews. A rush premium (7-day-inside bookings +25%) both monetizes urgency and shapes it (some customers book earlier to avoid it) — pure yield management your industry underuses. [§ 8.2 yield/fences, § 2.3, § 7]

Q2. Why do binding estimates lose to lowballs that blow up on move day — what does that teach the market?

A: It teaches customers that estimates are fiction and price-shopping pays — the lowballer wins the booking, then extracts on the truck. Your counter: educate at estimate time ("ask if it's binding and what triggers overages") and publish your no-surprise guarantee. Honesty needs marketing to compete with fraud. [§ 13 market dynamics, § 6.1 trust, § 2.3]

Q3. Does the packing upsell protect the customer or protect us from claims — what does packed-by claims data show?

A: Check it: owner-packed boxes typically show 3–5× the damage claims (and they're largely excluded from coverage). The upsell protects both — but lead with the customer's protection (their grandmother's china survives) and let the claims benefit be your quiet second win. [§ 6.2 failure prevention, § 2.3 framing, § 1.3]

Q4. Why do long-distance moves carry 3× margin and 5× complaints — which does the review profile show?

A: The complaints: long-haul moves involve agents, line-haul timing, and delivery windows you don't fully control — margin is real, but review damage compounds. Decide per lane: self-perform your core lanes, and either fix the partner-agent QA or exit the lanes where your reviews are purchased with margin. [§ 11.2 partner quality, § 6.1 reputation, § 1.3]

Q5. Why do customers rate the move by the last hour regardless of the first eight — what does crew scheduling do with that?

A: Peak-end rule: the unload's final hour (placement, reassembly, the walkthrough) writes the review. Crew scheduling usually front-loads energy and sloppens the end. Fix: reserve experienced crew energy for the finish, and script the final walkthrough as a designed moment. [§ 13 peak-end rule, § 8.2 blueprint, § 6.3]

Q6. Why do storage-in-transit customers become most profitable by accident — while we market the move, not the storage?

A: Because storage is the annuity (monthly, low-effort, high-margin) attached to a transaction you market. Make it deliberate: offer SIT proactively at estimate (closing-date gaps are common), price it visibly, and measure attach. The accident should be a product line. [§ 2.3 product design, § 1.3 annuity value, § 8.2]

Q7. Why do realtor referrals book at 70% vs 30% elsewhere — what does that relationship transfer?

A: Trust plus timing: the realtor's recommendation arrives at the decision moment with borrowed credibility. That transfer is worth money — formalize realtor relationships (response SLAs for their clients, closing-date guarantees) and track which realtors send the closers, not the shoppers. [§ 13 trust transfer, § 11.2 channel, § 1.3]

Q8. What breaks first at the fifth crew: the logistics system or the labor brokerage — which did we think we were scaling?

A: You thought you were scaling trucks; you're actually scaling labor quality and dispatch coordination. The fifth crew is where informal hiring ("guys who know guys") and whiteboard dispatch both fail. Staff the systems (real dispatch software, crew-lead development) before crew five. [A3 shifting bottleneck, § 1.1 structural, § 12]

Q9. Why do damage claims cluster on our most experienced crew?

A: Complacency or pace: veterans take shortcuts the manual warns against (their skill makes them fast, their speed skips padding steps) — or they get the hardest jobs (big houses, pianos). Stratify by job type first; if it holds within job type, it's complacency, and it's retrainable. [§ 6.3 stratification, § 13 overconfidence, § 4.4]

Q10. When summer ends, why do our best movers leave for warehouse jobs — what does the off-season offer them?

A: Nothing: seasonal layoff logic tells your best people to find year-round security, and warehouses offer it. The off-season offer could be: guaranteed minimum hours, training pay, inside maintenance/projects, or retention bonuses payable in spring. Your summer quality depends on your winter offer. [§ 2.1 seasonality, § 4.4 retention design, § 7 aggregate planning]

92. Courier & last-mile delivery

Q1. Does the same-day promise win business or convert every delay into a broken promise — what does the complaint log by promise-type show?

A: The log decides: if complaints cluster on same-day failures, your promise outruns your reliability — you're selling certainty you don't have. Either invest in the reliability (cutoff times, surge drivers) or re-promise (same-day by 8 PM vs. same-day period). A promise is a spec; meet it or rewrite it. [§ 0.3 specifiedBy, § 6.3 reliability, § 8.1]

Q2. Medical-specimen routes run perfectly while retail routes fail, with the same drivers. What does the medical route have?

A: Structure: fixed stops, trained contacts, chain-of-custody discipline, and consequences for failure. Retail routes have variability (new stops, absent recipients, vague instructions). Export the structure: stop databases with delivery notes, geofenced proof, and recipient protocols make retail behave more like medical. [§ 13 standard work, § 8.2 blueprint, § 12]

Q3. Why does the biggest client hold 35% of revenue and worst per-stop economics — which direction is it drifting?

A: Usually worse: anchor clients negotiate annual rate cuts while their stop density fragments your routes. Track the trend; if drifting, renegotiate with density data (offer volume incentives that improve YOUR economics — consolidated drops, off-peak windows) or diversify before the cliff.

[§ 11.4 concentration, § 10 route density, § 11.2]

Q4. Drivers' own vehicles outlast our fleet vehicles on identical routes. Which behaviors are we not measuring?

A: Ownership behaviors: gentle braking, warming up, reporting small issues early — owners protect their asset; employees drive a rental. Close the gap with telematics-scored driving (and bonuses for it) plus driver-assigned vehicles (pseudo-ownership changes behavior measurably).

[§ 13 ownership psychology, § 12 telematics, § 4.3]

Q5. When a gig app enters, why do we lose on-demand work but keep scheduled routes — which did we think was defensible?

A: You thought relationships defended everything; actually contracts defend scheduled routes while on-demand spot work is pure price/convenience competition — the app's home turf. Defend the contracted base (service levels, integration) and exit or niche the on-demand market (specialty items they can't handle). [§ 1.1 defensibility analysis, § 11.4, § 2.1]

Q6. What does the on-time-vs-margin curve say we deliver: speed, reliability, or price — which did the client sign for?

A: Align them: if you sell price and deliver speed-economics (rushed routes, high cost), margin dies; if you sell reliability at 95% on-time but staff for 99.9%, you're over-delivering unpaid quality.

Contract the metric, price it, and deliver exactly it — no more, no less. [§ 1.1 OrderWinner alignment, § 2.3, § 6.3]

Q7. Why do fuel surcharges lag fuel costs by a quarter, every quarter — what's the structural lag cost in a spike year?

A: The lag cost = spike amplitude × lag duration × volume — in a spike year, typically 1–2 margin points. The structure (quarterly surcharge resets) is the defect: move to monthly-indexed surcharges with published formulas (customers accept transparent indices). [§ 11.3 price fluctuation,

§ 0.3 governedBy, F1]

Q8. Why does per-stop cost fall with density until the driver's day exceeds legal hours — where is that point?

A: The point is computable: stops $\times$ service time + drive time vs. HOS limits. Density gains reverse past it (overtime, second driver, violations). Route-optimization should hard-cap at legal hours and model the marginal cost step at the cap — the "free" density past it isn't free. [§ 10 routing with constraints, § 14 HOS compliance, § 1.3]

Q9. Why do contracted routes churn at renewal despite zero service failures — what is procurement actually buying?

A: A savings story: procurement needs to show wins, and your zero-failure record gives them nothing to fix — so they buy the 5% cheaper bid to justify themselves. Preempt: offer structured savings yourself (efficiency commitments, rate locks) before they shop for them. [§ 13 procurement incentives, § 11.2, § 2.3]

Q10. When a client adds volume, why does our margin fall — which cost did pricing assume was fixed?

A: Route structure: pricing assumed new stops join existing routes (marginal cost ≈ 0), but volume growth forces new routes, vehicles, and drivers (step costs). Price volume tiers against step-cost reality: growth past a threshold should trigger repricing, not silent margin erosion. [§ 4.2 step costs, § 10, § 1.3]

93. Non-emergency medical transport (NEMT)

Q1. What does our hiring profile produce: a healthcare service with vehicles, or a transport company carrying patients — which does the training budget prove?

A: The training budget proves it: if training is driving-focused (defensive driving, routing), you're transport; if it includes patient handling, CPR, dementia care, HIPAA — healthcare. Patients and facilities experience whichever you trained. The healthcare identity wins facility contracts; fund it accordingly. [§ 4.4 training as identity, § 1.1 positioning, § 14]

Q2. Why do no-shows cluster by patient, not route — what would a patient-level flag change?

A: Everything: no-show behavior is patient-specific (dementia, ambivalence about dialysis, transport anxiety) and invisible at route level. A patient flag enables pre-calls the night before, facility coordination, and realistic scheduling — turning route-level chaos into managed exceptions. [§ 2.1 demand stratification, § 12 data, § 8.1]

Q3. Why do best-patient-feedback drivers have worst trips-per-hour — which does pay reward?

A: Trips-per-hour, presumably — paying for velocity while patients value the slow driver who helps them to the door. In NEMT, the care moment IS the service for your facility clients' satisfaction scores. Pay on composite metrics (on-time + complaint-free + assist compliance), not raw trips. [§ 13 Goodhart, § 6.3 SERVQUAL, § 1.3]

Q4. Dialysis patients are reliable volume concentrated in facilities we hold no contract with. What happens when one signs an exclusive?

A: Your volume vanishes overnight — you're riding facility-level demand with no facility-level relationship. Convert now: approach the facilities with performance data (on-time rates, patient satisfaction) and formalize before a competitor's exclusive does it for them. [§ 11.4 concentration/contract risk, § 11.2, § 2.1]

Q5. When we win a facility direct, broker volume from that facility drops immediately. What is the broker selling?

A: Access and administration — the broker's model is toll-collection on relationships you can own directly. The drop confirms direct contracts capture the margin. But note: brokers also smooth demand volatility; keep some broker mix as shock absorber while shifting the base direct. [§ 11.3 disintermediation, § 11.2 channel mix, § 11.4]

Q6. Why does adding vehicles reduce on-time performance for three months — which system is the real bottleneck?

A: Dispatch and driver-training: new vehicles arrive instantly, but new drivers take weeks to learn routes and patients, and dispatch complexity grows combinatorially. Capacity is never just the asset — it's the asset plus the trained human plus the coordination system. Sequence hiring before vehicles. [A3, § 4.4 training pipeline, § 7 dispatch]

Q7. When a facility discharges late, why does our wait cost land on us while trip rate stays fixed?

A: Because your contract prices the trip, not the wait — the facility's schedule slip is your unpaid inventory of driver-hours. Negotiate wait-time provisions (grace period + per-minute charge) or build expected wait into facility-specific rates. Unpriced externalities flow to whoever didn't contract them. [§ 11.2 contract design, § 11.3 externality, § 1.3]

Q8. On-time rates satisfy the broker contract while patient complaints tell another story. Which clock are we paid to serve?

A: The broker's clock (pickup within window) — while patients experience the whole journey (wait at return pickup, ride duration, driver manner). You're contract-optimized and experience-poor. The risk: facilities hear patient complaints and route around the broker. Measure both clocks. [§ 13 Goodhart, § 6.3, § 11.2]

Q9. Why do wheelchair vehicles — our scarcest asset — get assigned to ambulatory calls when dispatch is busy?

A: Because busy dispatch optimizes for "covered" not "correct": the wheelchair van is available NOW and the ambulatory patient fits. Each misassignment risks a later wheelchair call uncovered — your scarcest resource frittered on fungible demand. Lock asset-to-need matching in dispatch rules, with override logging. [§ 7 dispatch rules, A3 scarce resource, § 13]

Q10. Does our broker mix diversify risk or just diversify payment delays — what does days-to-pay by broker show?

A: The days-to-pay data answers it: if brokers vary 30–90 days, your "diversification" is a working-capital lottery. Diversify on payment behavior AND volume: drop or reprice the slow-payers (rate premium for slow terms is standard), and concentration risk shrinks as terms improve. [§ 4.1 working capital, § 11.2 terms management, § 11.4]

94. Limo & shuttle services**Q1. When TNC rates surge, why do we gain price-sensitive customers and lose the least price-sensitive — which does booking serve?**

A: Surge pushes deal-seekers to your flat rates while your best clients never price-shop — they're loyal to whoever answers reliably. If your booking experience (app friction, phone-only, slow confirmation) serves the patient deal-seeker over the impatient executive, you filter for the wrong segment. Fix booking speed; keep pricing flat. [§ 2.1 segment selection, § 8.2 blueprint, § 1.1]

Q2. Airport runs are steadiest volume and thinnest margins; weddings book six weekends a year. What is the growth plan chasing?

A: Check whether it chases glamour (weddings) at the expense of the annuity (airport/corporate). Wedding work is lumpy, seasonal, and coordination-heavy; airport contracts compound. Grow the base-load contracts deliberately; let weddings be margin cream, not the plan. [§ 1.3, § 8.2 base-load logic, § 1.1 focus]

Q3. Why is the cancellation policy enforced on leisure customers and waived for corporate — which cancels more?

A: Check the data: corporate accounts cancel more (their travelers' plans change constantly) and pay for the privilege via volume. That's fine — if priced. If the waiver is unpriced habit, you're donating flexibility. Build cancellation terms into corporate rates explicitly. [§ 8.2 fences, § 2.3, § 11.2]

Q4. Why do gratuity-included contracts produce happiest chauffeurs and most price-shopped accounts simultaneously?

A: Gratuity-included stabilizes chauffeur income (they stop performing for tips and deliver consistent service) but inflates the quoted price against competitors' tip-excluded quotes. The trade is real: keep it, and sell against it ("all-inclusive, no awkward envelope math") to accounts that value simplicity. [§ 13 incentives, § 2.3 price presentation, § 6.3]

Q5. Does fleet age matter to customers or only to the insurer — which vehicle earns most compliments per depreciation dollar?

A: Run compliments by unit: usually the answer is the newest sedan and the vintage/novelty vehicle, while mid-age units are depreciation without delight. Customers can't read model years; they read condition and cleanliness. Detail standards beat fleet youth for perception — insure and maintain, don't just replace. [§ 6.1 perceived quality, § 4.1 capital efficiency, § 13]

Q6. Why do affiliate farm-outs arrive at rates we'd never accept direct — and why do we accept them?

A: Because farm-outs fill idle capacity (any contribution > 0 at the margin) and keep chauffeurs busy. Accept strategically: farm-out work is fine as base-load filler, corrosive as growth strategy (the affiliate owns the customer). Cap affiliate share and track its true net rate. [§ 8.2 capacity yield, § 11.2 channel dependence, § 1.3]

Q7. Why does wedding pricing leave money in June while January corporate gets negotiated down — which season plans capacity?

A: June weddings are inelastic (date-fixed, emotional, booked long ahead) — raise them to market; January corporate is elastic and negotiates. Capacity should be planned around the contracted base-load (corporate/airport) with June priced as peak yield. Currently you're inverted: discounting the inelastic and pleasing the elastic. [§ 8.2 yield by elasticity, § 2.3, § 7]

Q8. Why do corporate accounts sign for sedans then book the party bus — two businesses run as one?

A: Because you sell "transportation" while clients buy occasions: the sedan decision (procurement, cost) and the party bus decision (the VP's offsite) happen in different brains. Segment the sale: contract the sedan utility, and market event vehicles separately to the same accounts' event planners. [§ 2.1 demand segmentation, § 1.1, § 2.3]

Q9. What does the fleet mix say: luxury brand with an airport problem, or transport utility in a tuxedo?

A: Count the hours: if most revenue-hours are airport/corporate sedan runs, you're a utility with luxury overhead (the stretch fleet's insurance and depreciation). Either sell the luxury inventory hard (events, weddings) or right-size it. Fleet mix is strategy in metal. [§ 1.1 identity, § 4.2 utilization, § 1.3]

Q10. Why do regular clients detect a chauffeur's departure before management — where does the service relationship live?

A: In the chauffeur: the client's continuity experience (routes known, preferences remembered) was person-bound. Institutionalize: client preference profiles in the dispatch system, and transition introductions when chauffeurs change. If clients know before you do, your feedback loop runs through your customers — unacceptable for a service business. [§ 13 relationship capital, § 12 CRM, § 6.3]

95. Childcare & daycare centers

Q1. Why do parents rate communication a strength while citing incident reports when leaving?

A: Because "communication" they rate is daily warmth (photos, pickup chat); incident reports are where communication becomes formal and alarming. The gap: incident handling that feels defensive or delayed poisons the whole record. Design incident communication (immediate call, written follow-up, prevention plan) as its own process. [§ 6.3 SERVQUAL assurance, § 8.2 recovery, § 13]

Q2. Why do waitlist families stay longer than immediate-enrollment families — what does scarcity do to commitment?

A: Scarcity filters and commits: waitlist families wanted YOU specifically (researched, waited), while instant-enrollees were solving an urgent problem and will solve the next one the same way. Guard the waitlist's integrity — it selects your best customers. [§ 13 commitment/scarcity, § 2.1 selection, § 2.3]

Q3. When a beloved teacher quits, why do her classroom's families follow within a quarter — whose brand did parents enroll in?

A: The teacher's: in childcare, the caregiver IS the product to the parent. Mitigate structurally: co-teaching (two attachment figures per room), center-level rituals parents bond to, and transition protocols when teachers leave. One-attachment classrooms are single-point failures. [§ 13 relationship capital, § 11.4 redundancy, § 6.3]

Q4. Why do infant rooms lose money at licensed ratios while preschool rooms carry the center — which does the expansion plan add?

A: If it adds infant rooms for pipeline logic ("feed the preschool"), it must price the infant loss as acquisition cost; if it adds rooms for their own economics, add preschool. The error is expanding infants without admitting they're a subsidy. [§ 1.3 cross-subsidy accounting, § 14 ratio constraints, § 1.1]

Q5. Staff call-outs cluster Mondays and Fridays. Which policy — PTO, scheduling, culture — have we never connected to the pattern?

A: The connection is usually PTO design: if taking a "real" day off costs a scarce PTO day, staff use sick calls for long weekends instead. Fix with floating holidays, weekend-adjacent shift swaps, and honest conversation — the pattern is a policy artifact, not a character flaw. [§ 13 incentives, § 4.4 scheduling design, § 14]

Q6. Does curriculum investment change enrollment or just justify tuition — which tours convert by which pitch?

A: Track tour conversion by pitch emphasis: usually the converting factors are warmth, cleanliness, and teacher stability (trust signals), while curriculum names justify the price afterward. Both matter — but spend on what the tour data says converts, and let curriculum be the retention story. [§ 6.1 perceived quality, § 2.3, § 13]

Q7. When we raise tuition, why do high-income families stay and two-income-stretched families leave — which does the waitlist replace?

A: The waitlist replaces leavers with whoever's next — often higher-income families who can pay, meaning raises gentrify your enrollment. If community mix matters to your mission (or subsidy compliance), the waitlist won't preserve it. Model tuition increases against your desired family mix, not just fill rate. [§ 2.1 selection effects, § 2.3, § 13]

Q8. Why do subsidy classrooms run at full capacity and negative margin — what strategic purpose do they serve?

A: Candidate purposes: license/accreditation requirements, mission, sibling-pipeline to private-pay rooms, or staff-utilization smoothing. If none apply, they're an unmanaged donation. Either price the purpose explicitly (it's worth X to us) or restructure the rooms. Negative margin with no purpose is just drift. [§ 1.1 deliberate strategy, § 1.3, § 14]

Q9. Families tour, enroll, and leave within 90 days at rates the waitlist hides. What happens in week six?

A: The honeymoon ends: week six is when routine friction (pickup lines, a teacher change, a biting incident, billing surprise) tests a relationship with no accumulated loyalty. The waitlist hides it because backfill is instant. Instrument a 90-day onboarding journey (check-ins at weeks 2/6/10) — churn prevention beats waitlist consumption. [§ 8.2 onboarding blueprint, § 2.2 leading indicators, § 13]

Q10. What does the financing assume: childcare business with real estate, or real estate play with a childcare tenant?

A: Read the debt structure: if the mortgage economics only work because the building appreciates, you're a real estate play running a childcare operator — which changes every decision (location over program, enrollment over quality). Misread identity leads to childcare decisions made by a landlord's logic. [§ 1.1 identity/structural, § 4.1 capital, § 13]

96. Tutoring & learning centers

Q1. Why do summer families not continue into the school year — what happens between August and September?

A: The context switch: summer buyers purchased childcare-with-enrichment or catch-up; September restores school (their primary structure) and tutoring feels redundant. Bridge in August: diagnostic results + a school-year plan presented as continuity ("keep the gains"), with scheduling fitted to school rhythms. [§ 2.1 demand seasonality, § 8.2 transition design, § 13]

Q2. SAT enrollments spike on counselor timelines, not our marketing calendar. Who owns the enrollment trigger?

A: The school counselor — your marketing is decoration on their advice cycle. Own the adjacency: counselor relationships, school presentations, and timing your campaigns to their recommendation windows (junior fall, test-date countdowns) rather than your fiscal calendar.

[§ 2.2 trigger mapping, § 2.3, § 11.2 channel]

Q3. Does the assessment predict outcomes or create a baseline that makes any progress visible — which do parents pay for?

A: Often the latter: a low baseline guarantees visible "growth" at re-test (regression to the mean plus any instruction helps). Parents pay for progress they can see. That's legitimate motivation design — as long as the assessment is honest. Keep assessments rigorous; the credibility compounds. [§ 13 measurement effects, § 6.1, § 14 compliance & standards]

Q4. Why do fast-improvers leave and plateaued students stay — which does the model depend on?

A: It depends on the plateaued: success completes the engagement ("done!"), struggle extends it. That's an awkward engine — your revenue needs slow progress. The ethical fix: package outcomes (score-guarantee programs, defined curricula with endpoints) so you profit from completion, and let referrals from successes replace the plateau annuity. [§ 13 incentive structure, § 2.3 packaging, § 6.1]

Q5. Why do unused hours cluster at package end — what does the redemption curve say about progress visibility?

A: The curve (front-loaded use, tail abandonment) says progress became invisible mid-package: early sessions show obvious wins, later sessions feel like maintenance with no visible delta. Counter with mid-package progress reports and milestone tests — keep the gain visible or the hours lapse. [§ 13 goal-gradient, § 6.1 perceived progress, § 8.2]

Q6. Why do group sessions carry better margins and worse outcomes than 1:1 — which does the website sell?

A: Probably group (the margin product) dressed in outcome language earned by 1:1 results. Honest architecture: sell 1:1 for measurable gaps, groups for enrichment/test-prep-cohort energy — and publish which fits what. Mismatched selling churns parents at results time. [§ 1.1 product-positioning fit, § 6.1 honesty, § 2.3]

Q7. When grades improve, why do parents credit the school and cancel us — whose visibility failure is that?

A: Yours: the improvement is public at school, private at your center, and nobody connects the two. Fix attribution visibility: parent reports mapping tutoring content to the improved units, teacher-comment tracking, before/after work samples. Unclaimed credit doesn't renew. [§ 13 attribution, § 6.1, § 2.3 retention]

Q8. Our best tutors build waitlists then leave to tutor independently, taking the list. What does the hiring conversation say about that risk?

A: Nothing, probably — and it should address it directly: non-solicitation terms (where enforceable), plus positive retention (revenue share on their waitlist, lead-tutor tracks). You can't prevent independence; you can make staying richer than leaving. [§ 13 incentives, § 11.4 key-person risk, § 1.1]

Q9. Why do referrals come from parents whose kids graduated two years ago — what does the lag do to pipeline math?

A: The lag means today's pipeline reflects service quality from two years ago — improvement efforts take two years to show in referrals. Pipeline math must discount current marketing accordingly, and you must survive the lag: the referral engine runs on a delay, so don't cut quality in a quiet quarter. [§ 2.2 lag structure, F1 time-indexing, § 13]

Q10. Which client churns first when money tightens: the deficit-remediator or the anxiety-insurer — what does that reveal?

A: The anxiety-insurer churns first (the anxiety was discretionary); the deficit-remediator stays (visible need, measurable stakes). Your revenue mix by motivation type determines your recession resilience — know the mix, and market the necessity framing when household budgets tighten. [§ 2.1 demand elasticity by motivation, § 11.4, § 13]

97. Self-storage facilities

Q1. When a competitor opens a mile away, why do we lose new rentals but keep every tenant — which number does revenue management watch?

A: It should watch move-in rate (the bleeding), not occupancy (the illusion): existing tenants are inertia-locked (moving stuff is misery), so occupancy holds while your demand pipeline dies. Counter on acquisition (rate-match for new move-ins, move-in specials) — your base defends itself. [§ 13 switching cost/inertia, § 2.2 leading indicators, § 8.2]

Q2. Manager conversion varies fourfold on identical phone inquiries. Which script variable does mystery-shop data isolate — why isn't it training?

A: Usually: quoting price-first vs. needs-first ("what are you storing?"), plus offering to hold a unit. The mystery-shop data has the answer; converting it to training is a standard-work exercise with a 4× payoff. That it isn't training yet is the actual finding. [§ 13 standard work, § 6.3, § 2.3]

Q3. At 90% occupancy, why do we discount to fill the last 10% — what does the rate-vs-occupancy curve maximize?

A: Nothing — discounting at high occupancy undercuts your own scarcity value. Revenue management in storage: RAISE street rates above ~85% occupancy (scarcity pricing), and let occupancy sit at 92% at higher rates rather than 98% at discounted ones. The curve's optimum is rate, not fullness. [§ 8.2 yield management, § 2.3, § 13]

Q4. Why do promo tenants stay longer than full-price tenants — inverting the promo's assumption?

A: Because promo tenants are planners (they hunted a deal = organized, long-need storage) while full-price walk-ins are crisis renters (divorce, eviction) whose storage need evaporates. The promo selects for duration. Re-examine: promos attract the customers you actually want. [§ 2.1 selection effects, § 13, § 2.3]

Q5. Does the auction process recover losses or just process them — what's net recovery per delinquent unit after labor?

A: Usually processing: auction proceeds minus labor, notices, and legal steps often nets near zero. The real lever is earlier: autopay penetration, early-delinquency contact (day 5 call beats day 45 letter), and overlook timing. Prevention outperforms liquidation 10:1. [§ 6.2 prevention vs failure cost, § 8.2, § 1.3]

Q6. Why do climate-controlled units command 40% premiums with identical delinquency — which product is the real business?

A: Delinquency parity says payment behavior is unit-type-independent, so the climate premium is nearly pure margin. But watch demand: climate units serve different contents (business records, furniture vs. overflow junk). If climate occupancy holds, expand climate at conversion — the premium is the product. [§ 1.3 margin structure, § 2.1, § 10]

Q7. Why do business tenants pay reliably and refer constantly at 10% of marketing spend?

A: Because business storage solves operational problems (inventory, tools, files) with real ROI — sticky and talkative in contractor networks — while marketing targets residential life-events. Rebalance: contractor yards, e-commerce seller outreach, and trade-relationships. Business tenants are your best segment hiding in plain sight. [§ 2.1 segment economics, § 2.3, § 1.3]

Q8. Longest-tenured tenants pay least per square foot. Which number does the annual-increase policy move?

A: If increases apply uniformly, they move your newest tenants toward churn while the legacy base stays cheap. Storage pricing convention (aggressive existing-customer rate increases) works because moving is misery — but apply increases by tenure curve: catch the legacy base up gradually. [§ 8.2 ECRC logic, § 13 inertia economics, § 2.3]

Q9. Are we renting space or renting procrastination — which does length-of-stay data say people buy?

A: Procrastination: stay lengths far exceed move-in intentions ("one month" becomes 19). The product is deferred decisions. Implication: don't oversell move-out support; do optimize for move-in ease and autopay — the business model runs on benign inertia. [§ 13 behavioral ops, § 2.1, § 14 triple bottom line (people)]

Q10. Why does the facility sit next to its highest use case (apartment turnover) while marketing targets everyone?

A: Because marketing is generic while demand is hyper-local: apartment-complex turnovers within 2 miles are your highest-conversion demand. Geo-fence the marketing: apartment-complex partnerships, move-in-week flyers, "first month" timing matched to local lease cycles. [§ 10 location demand, § 2.3 targeting, § 13]

98. Dry cleaners & laundromats

Q1. When we raise prices, why do garment counts drop before customer counts — what are customers quietly taking elsewhere?

A: The marginal garments: they keep bringing suits and dresses (must-clean) and start home-washing or wearing-longer the casual items (shirts, knits). You're losing volume, not customers — meaning the price increase hit elastic demand. Consider category-tiered pricing instead of blanket increases. [§ 2.1 elasticity by category, § 2.3, § 13]

Q2. If casual dress codes are permanent, why is the counter garment mix unchanged from a decade ago — which categories still grow?

A: The mix should shift to what's left: household items (comforters, drapes), specialty (leathers, wedding preservation), and wash-and-fold (time-poor professionals). The shirt-press volume isn't returning. Retool counter capacity and marketing toward the growing categories. [§ 2.1 demand trend, § 1.1 adaptation, § 1.3]

Q3. Does eco-solvent positioning attract customers or reassure existing ones — which new customer ever cited it?

A: Check intake: if new customers never cite it, it's retention reassurance, not acquisition. That's still worth something (defends against the green-branded competitor) — but stop spending acquisition money on it; put it in retention comms instead. [§ 2.3, § 6.1, § 13 attribution honesty]

Q4. Why do wash-and-fold customers stay for years while dry-cleaning customers defect over one damaged blouse?

A: W&F is a logistics habit (weekly rhythm, low emotional stakes); dry cleaning is trust-custody of valued garments — one failure is a betrayal, not an error. Price and process accordingly: garment-care customers need claims handled instantly and generously; the recovery defines the relationship. [§ 8.2 service recovery, § 6.1, § 13 habit vs trust]

Q5. Why do card machines out-earn coin per square foot while coin customers buy the vending and drop-off?

A: Two customer systems: card users are convenience-driven (higher income, machine-only); coin users are cash-economy regulars who linger and buy the ancillaries. Serving both is fine — but don't let card-machine economics talk you into removing the coin customers' reasons to linger. [§ 2.1 segmentation, § 10 layout, § 1.3]

Q6. Where do capex decisions point: garment-care business with a laundromat, or laundromat with a counter?

A: Follow the capex: if machines dominate spending, you're a laundromat operator with a counter habit — and counter decline will continue unmanaged. Decide: garment care needs different capex (pressing, delivery vans, POS); laundromat needs machine refresh and amenities. One identity gets the money. [§ 1.1 identity, § 4.1 capex as strategy, § 13]

Q7. Why do commercial accounts pay at 45 days and churn at 18 months — which does pricing assume?

A: Pricing probably assumes prompt pay and longevity — both wrong. Commercial linen work is price-shopped annually and pays slow. Price it with the terms included (payment-term surcharges or prepay discounts) and treat 18 months as the expected life, making acquisition cost match. [§ 4.1 working capital, § 2.1 churn reality, § 1.3]

Q8. Why do busiest hours coincide with worst machine uptime — which machine fails first and why?

A: Peak load exposes the weakest machine: high utilization accelerates wear, and the failure then removes capacity exactly when demand peaks (Kingman for washers). Track failure by machine age/model; preemptively replace the statistical worst, and schedule PM in trough hours. [§ 4.3 reliability/maintenance timing, A4, § 6.3]

Q9. The delivery route acquires customers who never visit the store. What happens when the driver changes?

A: Churn spikes: the driver is the only human face of a faceless service — their knock, their reliability, their texts. When they change, customers with no store attachment drift. Systematize the relationship: route communication from the brand (texts, notes), not the driver's personal phone. [§ 13 relationship institutionalization, § 8.2, § 11.4 key-person]

Q10. Why does attendant-on-duty cost more in payroll than it saves in shrink — has the test been run?

A: Run it: go attendant-free in a measured window (cameras + remote monitoring) and compare total cost (payroll saved vs. shrink + machine abuse + cleanliness labor). Many operators find attendants pay for themselves in machine longevity and drop-off sales, not shrink — measure the right ledger. [§ 1.3 full-cost analysis, § 14, § 13]

99. Private security companies

Q1. Does the technology stack compete with guard revenue or extend it — which does sales pitch?

A: Sales pitches guards (the revenue they know) while tech (remote monitoring) offers better margins at lower price — cannibalizing their own book. Resolve deliberately: position tech as the force multiplier (guards + cameras + monitoring = integrated service) with comp plans that pay on the bundle. [§ 12 digital layer, § 13 sales incentives, § 1.1]

Q2. When a client has an incident, they increase our hours rather than fire us. What does the contract insure?

A: Their anxiety and their liability posture — the incident proves the threat was real, and more coverage is the response. You sell risk-transfer and demonstrable diligence, not incident prevention per se. Document that value (deterrence reporting, incident response quality) — it's what renews contracts. [§ 11.4 risk-transfer product, § 6.1, § 2.3]

Q3. Why do guard-wage increases get absorbed at exactly the accounts whose contracts allow escalation?

A: Because invoking escalation clauses feels relationship-risky, so account managers absorb cost to protect rapport — donating margin with no credit received. Make escalation invocation routine (annual, indexed, automatic-notice) so it's process, not confrontation. [§ 0.3 governedBy, § 11.3, § 13 conflict avoidance]

Q4. What does the client CFO believe they bought: labor with uniforms or risk management with guards — which does renewal pricing assume?

A: If pricing assumes labor (cost-plus on wages), the CFO shops you against staffing agencies. Risk-management framing (incident data, insurance implications, compliance documentation) supports value pricing. The renewal conversation's framing determines your ceiling. [§ 1.1 positioning, § 2.3 framing, § 6.1]

Q5. Clients renew while complaining about guard quality. What are they buying that we never measure?

A: Continuity and non-event: no incidents, no vendor-switching hassle, no explanation to their board. The complaints are negotiating posture, not exit intent. What you never measure: incidents-that-didn't-happen and switching-cost relief. Measure presence/consistency metrics and price the quiet. [§ 6.1 prevention value, § 13, § 1.3]

Q6. Our best officers leave for police departments exactly when site knowledge peaks. Which accounts absorb that loss?

A: The accounts those officers served — losing trained judgment about the site's patterns. Structural responses: document site knowledge (post orders as living documents), pay-progression that narrows the police-gap at year 2–3, and recruit from the pipeline that police departments reject but clients accept. [§ 13 knowledge codification, § 11.4 talent pipeline, § 4.4]

Q7. Overnight posts are hardest to staff and carry the least quality-willing clients. What would repricing nights do?

A: It would reveal which clients actually need nights: repricing to true cost (differential wages + supervision difficulty) pushes marginal clients to camera-only solutions (which you can sell) and keeps committed ones at sustainable rates. Nights priced like days is a subsidy to clients who value neither. [§ 2.3 cost-true pricing, § 8.2, § 4.4]

Q8. Why do mobile patrol accounts carry better margins and worse incident documentation than posted accounts?

A: Because patrol is report-what-you-saw with no one watching, while posted officers have client contact daily (documentation demanded). Patrol's margin comes partly from its invisibility. Instrument patrol (GPS-verified checkpoints, photo logs) — the documentation gap is a liability and a sales asset waiting to be productized. [§ 12 verification tech, § 6.3, § 14 liability]

Q9. Why do field supervisors visit accounts that complain least — what does visit-log vs incident-log overlap show?

A: It shows supervision follows comfort, not risk: quiet accounts get visits, troubled accounts get avoidance. Invert the routing: visit frequency keyed to incident/complaint leading indicators. Supervision allocated by complaint-avoidance is how contracts get lost "unexpectedly." [§ 13 avoidance behavior, § 2.2 leading indicators, § 6.3]

Q10. When we lose a contract, why does the stated reason (price) never match exit-interview truth (the supervisor) — which do we fix?

A: The supervisor — but price is the polite exit story. Fixing supervision (visit quality, guard development, client contact) is the real retention lever; price-fixing responds to the lie. Require two-source exit data (client contact + their on-site staff) before accepting "price" as cause. [§ 13 face-saving exits, § 6.3, § 2.3]

100. Pet grooming & boarding

Q1. Why do daycare owners churn when the dog "graduates" to staying home calmly — the outcome we sold?

A: Because you sold the outcome (calm dog) while selling the subscription (daycare) — success terminates the need. Reposition what daycare sells: socialization maintenance and the owner's lifestyle freedom (not just energy-burn). Graduates should convert to 1–2 day/week "social memberships," not zero. [§ 13 goal-completion churn, § 2.3 product reframing, § 8.2]

Q2. Why do intake requirements filter out easy revenue while cleared clients generate incident reports?

A: Because requirements screen paperwork (vaccines, tests) not behavior-fit: a documented dog can still be a chaos agent, while an undocumented calm dog is easy revenue lost. Add behavioral observation (first-day evaluation scoring) to the gate, and re-check incumbent dogs annually. [§ 6.3 acceptance criteria design, § 14 safety, § 13]

Q3. Why do boarding regulars' daycare packages erase the daycare margin?

A: Because packages were priced against casual daycare rates, not against power-user consumption: your best customers buy the package and use it at frequencies that invert the per-day economics. Price packages with usage caps or tiers sized to actual power-user behavior. [§ 8.2 subscription economics, § 2.3, § 1.3]

Q4. Why did cat boarding fill instantly then plateau at 40% — which assumption about cat owners was wrong?

A: The assumption that cat owners travel-board like dog owners: cats hate relocation, and most cat owners prefer in-home sitters. Your initial fill was the pent-up minority (medical cats, long trips). The plateau is the true market size — capacity beyond it was over-built for cats. [§ 2.1 demand estimation error, § 4.2 capacity, § 13]

Q5. What survives WFH: the pet care we provide, or the guilt relief they were buying?

A: The guilt relief shrinks (the dog is no longer alone), but what survives is socialization needs and owner convenience (meetings, focus time). If WFH halves your demand, reposition to socialization packs and flexible half-days — the dog's social needs outlive the owner's commute. [§ 2.1 demand driver analysis, § 2.3 repositioning, § 13]

Q6. Why do difficult-dog surcharges get waived for exactly the dogs costing most labor — whose decision is that?

A: Front-line staff waiving under pressure (the apologetic owner, the scene at pickup) — a policy without enforcement mechanics. Make surcharges system-driven (flagged in booking, auto-applied) so waiving requires a manager override. Compassionate exceptions then become deliberate, not default. [§ 0.3 governedBy enforcement, § 13, § 1.3]

Q7. Why do a departing groomer's doodle clients follow her while short-hair clients stay — what does that split reveal?

A: That high-skill, relationship-heavy grooms (doodles: breed-specific cuts, anxious dogs, trusting owners) are person-bound, while commodity grooms (short-hair baths) are shop-bound. The shop owns the commodity; the groomer owns the craft relationship. Build shop-level doodle loyalty (consistent pricing, style notes, rebooking systems) before the next departure. [§ 13 person vs institutional loyalty, § 11.4, § 12 style records]

Q8. Does the report-card/photo service reduce anxiety or create complaint-generating expectations — which reviews mention it?

A: Check: if complaints cite photo discrepancies ("looked stressed in the picture"), the service creates evidence that invites scrutiny; if reviews praise it, it's your differentiator. Usually both — meaning the service works but must be curated (send the good moments, never the kennel-sulk photo). [§ 6.3 assurance, § 13 expectation management, § 8.2]

Q9. Why do holiday bookings fill by October while staffing assumes December availability, every year?

A: Because booking lead-time data never reached the staffing calendar: holiday demand is locked by October, but hiring starts in November. Move the recruiting trigger to the booking curve (when holiday occupancy hits 50%, start seasonal hiring) — the data to fix this exists in your own reservation system. [§ 2.2 leading indicators, § 7 capacity planning, § 4.4]

Q10. Grooming books 6 weeks out; boarding books 6 days out. Which does front-desk capacity protect?

A: Usually boarding (the phones ring for boarding emergencies), starving grooming's long-horizon scheduling (which needs proactive rebooking calls). Protect both: groom rebooking should happen at pickup (book the next groom before they leave), while the desk handles boarding's short-fuse demand. [§ 8.2 demand horizon management, § 2.3, § 7]

Appendix — Reference Key

Ontology sections cited throughout (Part I of this book):

- **Part 0** — Formal foundations: continuants/occurents, Events, States; F1 (time-indexing), F2 (dual role); D1–D5 (disjointness/cardinality); core relations ([measuredBy](#), [governedBy](#), [specifiedBy](#), [constrainedBy](#)...)
- **Part 1** — Strategy layer: CompetitivePriority, OrderQualifier/OrderWinner (Hill), Structural vs Infrastructural decisions, Product–Process Matrix, Focus, Sand Cone, Decoupling Point, Learning Curve
- **Part 2** — Demand layer: independent/dependent demand, patterns (seasonality, cycles), forecasting methods, demand management (influencing/accommodating)
- **Part 3** — Process core: Process/Activity/FlowUnit, VA/NVA/NNVA, process states (Idle/Blocked/Starved/Down)
- **Part 4** — Resources: capacity (design/effective/utilization), OEE, Maintenance policies (preventive/predictive/RCM/TPM), bathtub curve, reliability math, job design, work measurement, skill matrices
- **Part 5** — Inventory: classes, EOQ/Newsvendor/(Q,r)/ABC, risk pooling, square-root law
- **Part 6** — Quality: Garvin dimensions, Cost of Quality, SPC (common/special cause), poka-yoke, QFD, SERVQUAL gaps
- **Part 7** — Planning & control hierarchy: S&OP, MPS/MRP, dispatch rules, pull systems
- **Part 8** — Service operations: Kendall queues, Erlang C, Kingman VUT, blueprinting, service recovery paradox, revenue/yield management (fences, overbooking)
- **Part 9** — Project management: CPM/PERT, critical chain, EVM
- **Part 10** — Facility & layout: layout types, line balancing, location/routing models
- **Part 11** — Supply chain: SCOR, Kraljic, sourcing, bullwhip, coordination failure, risk & resilience (TTR/TTS, BCP)
- **Part 12** — Digital ops: sensing/platform/model/intelligence layers

- **Part 13** — Human & behavioral ops: Goodhart's law, behavioral biases, incentives, standard work, knowledge management
- **Part 14** — Sustainability & safety: leading/lagging safety indicators, compliance
- **Part 15** — Axioms: **A1** Little's Law • **A2** $\text{Throughput} \leq \min(\text{bottleneck}, \text{demand})$ • **A3** Bottleneck set (time-varying) • **A4** Variability–buffer law • **A5** Conservation of flow • **A6** OEE decomposition • **A7** Risk pooling • **A8** Traceability • **A9** Learning curve • **A10** Strategy–structure consistency

Method note: each answer applies the ontology as an analytical lens to the question's mechanism. Answers are diagnostic guidance, not empirical claims about any specific business; where the ontology implies a test ("stratify by crew", "compute the cohort"), the answer says what to measure rather than asserting the result.

Citation convention: in each bracket, the [§ n / An / Dn / Fn](#) key is a reference to the ontology in Part I (Part/Section or Axiom). The short descriptor after the key is an **application note** — it names how the ontology class is being applied to the question (e.g., [§ 11.4 key-person risk](#) applies Part 11.4 Risk & Resilience to a key-person exposure). Descriptors are interpretive labels, not quoted ontology terms; verbatim ontology terms are reserved for the class names themselves (Kingman VUT, OrderWinner, Goodhart's law, bathtub curve, etc.).