

Standardize the record, not the reader: Decision records for consumer AI agent markets

Erik Postnieks

Center for Decision Accounting

erik@decisionaccounting.org

<https://decisionaccounting.org/>

Project page and version-matched replication package: <https://decisionaccounting.org/calvano-drcs/> Direct replication archive: <https://decisionaccounting.org/research/calvano-drcs/calvano-drcs-v1.2-replication-package.zip>

Working paper v1.2

July 28, 2026 revision

© 2026 Erik Postnieks. All rights reserved.

Abstract

Consumer AI-agent markets turn on loyalty: whether an agent works for the user or is steered by vendors, marketplace incentives, merchant fees, or the provider's own inventory. A successful purchase does not reveal whether the agent rejected a better seller, lower total price, safer payment path, or better return policy. This paper asks what an agent should record before acting and when that evidence changes conduct.

I adapt the two-firm Q-learning pricing laboratory of Calvano, Calzolari, Denicolò, and Pastorello and attach structured decision evidence to each pricing choice. The experiment separates evidence production, forecast scoring, a changed private objective, outcome-contingent restrictions, dose-matched random interruption, rejected-alternatives review, and a proposed several-reader design. The corrected implementation uses the published price grid, imposes the stated restriction duration, and runs the expensive simulation loop with compiled Numba kernels.

Producing the evidence path without changing rewards or oversight leaves pricing unchanged. Forecast scoring also leaves pricing unchanged in this implementation. Outcome-contingent restrictions can change conduct, and at the higher audit intensities tested their targeting matters relative to an equally frequent random interruption. The rejected-alternatives and several-reader additions produce no incremental effect because their implemented triggers do not activate. These nulls define the result: a record is evidence, not an incentive. It changes conduct only when an operative reader or rule can make the evidence consequential. The proposed Decision Record & Control Standard Consumer E-Commerce profile (DRCS-EC) translates that experimental lesson into fields for selected and rejected options while leaving its field performance for separate testing.

Keywords: consumer AI agents; agent marketplaces; algorithmic collusion; decision records; audit; reinforcement learning; antitrust; system welfare.

JEL codes: C63, C72, D43, D83, K21, L13, L41, L51.

Highlights.

- Producing decision evidence alone leaves learned pricing unchanged.
- Forecast scoring improves no pricing outcome in this implementation.
- Targeted restrictions outperform matched random interruption at high intensity.
- The rejected-alternatives trigger is inactive and adds no effect.
- Reader diversity still requires genuinely distinct active reader rules.

Plain language summary

An agent can produce an acceptable result while quietly passing over a better option for the user. A receipt shows what the agent chose; it does not show the options the agent considered, rejected, and predicted. The experiment records that missing decision evidence inside a Q-learning pricing laboratory. Evidence production by itself changes nothing. A consequence tied to harmful outcomes can change prices, while an alternatives test that never fires cannot. The result separates three jobs that are often bundled together: creating evidence, reading it, and attaching a consequence. The proposed standard supplies a common record so regulators, marketplaces, courts, firms, and users can test those jobs separately.

Policy ask

NIST and the European Commission should consider a DRCS-EC pilot standard: a testable consumer-commerce decision-record profile that would let regulators and researchers evaluate whether an agent’s stated loyalty constraints match the options it considered and rejected. The pilot should specify fields needed to test that question: user preferences, sellers and products considered, total price, shipping terms, return terms, payment path, checkout friction, seller reliability, selected option, rejected options, and a declared system-welfare effect. Those fields can also reveal candidate steering patterns—repeated routing into one marketplace, affiliated inventory, preferred sellers, sponsored placements, or data channels that strengthen a gatekeeper. NIST can separately review DRCS Core 1.0 as the broader framework that EC adapts. Ranking algorithms, reader institutions, and consequences remain choices for marketplaces, rating services, enterprise buyers, regulators, courts, and users.

Summary of findings

- Activating the evidence-computation path leaves learned pricing exactly unchanged under common seeds.
- Forecast scoring does not change pricing; calibration improves during learning in both the scored and unscored evidence paths.
- Putting consumer surplus into reward changes the learner’s private objective and sharply changes conduct.
- Outcome-contingent restrictions reduce supracompetitive pricing, with a nonmonotonic response to audit intensity.
- Targeting adds little at low intensity and a large effect at the two highest intensities tested.
- The rejected-alternatives add-on produces no incremental effect because its trigger is effectively inactive.
- The several-reader add-on also produces no incremental effect because its composite never activates.
- Outcome-only ratings cannot show whether an agent worked for the user or was steered by third-party incentives.
- Outcome-only ratings cannot show whether agent markets are reinforcing platform monopoly through hidden steering.
- DRCS-EC proposes evidence fields for evaluator agents, including rejected alternatives and declared system-welfare effects.
- A pilot should test a common record layer with heterogeneous readers and consequences before any performance claim.

Contributions

1. The paper converts the consumer-agent marketplace problem into a record-standard problem: outcome-only ratings cannot show whether an agent served the user or third-party interests.
2. It specifies a corrected replication design for testing DRCS-EC inside the Calvano algorithmic-collusion environment, with DRCS Core 1.0 as the broader framework EC adapts.
3. It separates the record-evidence computation path, forecast scoring, welfare-in-reward, outcome-contingent enforcement, random interruption, and a record-specific rejected-alternatives trigger.
4. It estimates the dose response and compares outcome-contingent restrictions with random interruption matched on realized restriction exposure.
5. It reports two design failures: the rejected-alternatives trigger is effectively inactive and the several-reader composite never activates.
6. It identifies rejected alternatives and system welfare as candidate EC fields whose value requires active triggers and field validation.

- It offers a Prat-motivated institutional design proposal: standardize the record, preserve reader diversity, and test how market and legal readers use the evidence.

Reader's guide

The main paper states the problem, experimental results, and proposed standard. Readers who want the takeaway can read Sections 1-5. Referees and replication readers can use Appendices A-D for the simulation specification, robustness results, replication protocol, and schema documentation. The project page supplies the version-matched code, results, figures, and checksums.

How this paper fits the literature

Prior work	Established result	This paper's bounded addition	What this paper does not establish
Calvano and coauthors (2020)	Independent Q-learning pricing agents in a differentiated duopoly can learn supracompetitive prices under the authors' conditions.	Uses the corrected published environment and treats oversight design as the experimental variable.	Field prevalence, a new theory of collusion, or a universal enforcement design.
Calvano and coauthors (2023)	Supracompetitive prices may be genuine or spurious; the rate and mode of exploration help determine whether learned conduct is supported by punishment strategies.	Keeps the exploration design explicit and limits the enforcement claim to the tested learner and environment.	That every high-price outcome is collusion or that an oversight treatment generalizes across exploration designs.
Klein (2021); Calvano and coauthors (2021)	Learned supracompetitive pricing can survive sequential moves and imperfect monitoring.	Holds the market environment fixed while varying incentives, outcome-contingent enforcement, audit intensity, and a record-specific add-on.	That every pricing technology or monitoring regime produces the same result.
den Boer, Meylahn, and Schinkel (2026)	Q-learning collusion can require commercially irrelevant timescales and synchronized algorithm choices, limiting practical relevance.	Uses the Calvano laboratory as a controlled enforcement test bed and makes no prevalence or deployment claim.	That the laboratory establishes a current cartel risk in deployed pricing markets.
Schildknecht (2026)	An open Python implementation replicates the Calvano study's principal pricing and punishment results.	Positions this paper as an oversight-treatment extension after the corrected-grid rerun.	Novelty from reproducing the baseline environment itself.
Assad et al. (2024)	Adoption of algorithmic pricing was associated with higher gasoline margins, especially when rivals also adopted.	Supplies a controlled environment in which the evidence attached to each decision is observable.	A causal estimate for deployed consumer-agent markets.
Asker, Fershtman, and Pakes (2022, 2024); Abada and Lambin (2023)	What an algorithm evaluates and how it explores can determine whether prices become supracompetitive.	Keeps the learner fixed and gives counterfactual decision evidence to an outside evaluator.	That records dominate learner redesign or changes to exploration.
Brown and MacKay (2023)	Pricing-technology asymmetries can raise prices without collusive learning.	Shows why observed prices alone may not identify the mechanism and tests an additional decision-evidence layer.	That a decision record alone identifies every source of high prices.
Harrington (2018); Harrington and Imhof (2022)	Legal analysis and outcome-based statistical screens can identify suspicious market patterns.	Adds decision-level evidence about what an agent saw, chose, predicted, and rejected.	A substitute for market-wide screens, discovery, or legal proof requirements.
Hartline, Long, and Zhang (2024); Hartline, Wang, and Zhang (2025)	Sellers can retain data that supports statistical audits of whether a pricing algorithm satisfies a non-collusion test.	Separates record-evidence computation, audit intensity, consequence dose, and a record-specific incremental trigger.	That this paper's audit trigger is the optimal legal or statistical test.
Chan et al. (2024)	Agent identifiers, activity logs, real-time monitoring, and related infrastructure can improve visibility into AI agents.	Specifies rejected alternatives and declared welfare effects as candidate fields and defines an experiment for their incremental effect.	A complete infrastructure or privacy design for deployed agents.
European Union AI Act, Article 12	High-risk AI systems must automatically log events sufficient for traceability and post-market monitoring.	Proposes record content and a pilot evaluation program beyond the statutory event-logging baseline.	That DRCS-EC is legally required or sufficient under the Act.
Prat (2005)	Transparency about actions can induce conformism toward a known evaluator.	Motivates a common record with distinct active reader rules as a future experiment.	That heterogeneous readers eliminate conformism or changed behavior in the archived implementation.

Prior work	Established result	This paper's bounded addition	What this paper does not establish
Aguirre et al. (2020); Kolt (forthcoming)	AI assistants raise a loyalty problem and may require fiduciary-style governance.	Defines decision evidence that could make a loyalty claim testable and subjects part of that design to a controlled experiment.	A completed legal duty, conformance regime, or field test of consumer-agent loyalty.

This paper's contribution is narrower than the general proposal to govern agents through records. Existing scholarship already proposes retained records and empowered audits, European law already requires event logging for some high-risk systems, and Schildknecht (2026) supplies an open Python replication of the baseline environment. Calvano and coauthors (2023) show why exploration design and punishment-supported conduct must be distinguished from a high-price outcome. den Boer, Meylahn, and Schinkel (2026) question the commercial relevance of collusion that depends on long learning horizons and synchronized algorithm choices. Those results make the laboratory useful here as a controlled enforcement test bed, without treating it as evidence of prevalence in deployed markets.

The corrected experiment separates record-evidence computation, forecast scoring, welfare-in-reward, outcome-contingent restrictions, dose-matched random interruption, and a rejected-alternatives add-on. The outcome trigger uses realized profit and consumer surplus and does not require the structured record. The alternatives add-on therefore supplies the record-specific incremental comparison. DRCS-EC translates the experimental objects into a proposed field specification. The experiment remains a two-agent Q-learning laboratory with truthful machine-state evidence, complete sample records materialized at session end, and a known welfare measure; it does not establish durable per-decision filing or the performance of the proposed standard in deployed markets.

The table positions a design contribution. The three selected examples in this paper—seller-side pricing agents, consumer-agent steering, and platform concentration—illustrate different uses of a record. They do not establish one common causal mechanism or prove that every deployed system has the same failure mode.

1. The policy problem

AI-agent marketplaces create a loyalty problem. A consumer agent is supposed to work for the end user. It can also be pulled toward the vendor, the platform, a merchant paying a fee, or the agent provider's own inventory and partnerships. The user may never see that pull. The purchase can arrive on time, the price can look acceptable, and the user can leave a good rating while the agent quietly rejected a better option.

The practical question is simple: what would a reviewer need to see to test that claim? A receipt shows the selected item. A decision record can also show what the agent predicted, which alternatives it considered, which it rejected, why it recommended one option, and what happened afterward. The record is evidence; it changes conduct only when a payoff, rule, or authorized reader can use it.

That loyalty problem can also create a platform-power concern. A marketplace that controls agent distribution, rankings, data, and checkout can steer users toward favored sellers or affiliated services while the receipt still looks ordinary. This is a selected illustration of a potential gatekeeper mechanism, not an estimate of concentration or capture in deployed consumer-agent markets.

That is the central blind spot. A marketplace can show the receipt. It cannot prove, from the receipt alone, that the agent served the user rather than a third party. Outcomes and user ratings are too thin because they see only the transaction the user observed. The missing evidence is the decision itself: the options considered, the option selected, the options rejected, and the broader effect of the choice.

Prat's conformism result motivates a consumer-agent design question. If one known evaluator controls the rating, builders can optimize for that evaluator. DRCS-EC proposes a common record that several institutions could assess with different rules and powers. That reader-diversity proposal has not yet been identified experimentally

here: the archived several-reader score added no active consequences and produced no incremental behavioral effect. Laboratory and field tests with genuinely distinct reader rules remain necessary.

The proposal joins decision evidence to consequence evidence: what the agent saw, what it chose, what it rejected, what happened afterward, and which institution may act. Marketplaces, user assessment agents, enterprise buyers, regulators, courts, and competing rating services could apply different tests to the same record. Rank, placement, procurement access, certification, enforcement, and proof are possible institutional consequences whose effects require separate evaluation.

The same problem appears in the Calvano study. The authors showed that reinforcement-learning pricing agents can learn supracompetitive prices without a collusive instruction or message between them. Later work distinguishes punishment-supported collusion from high prices created by exploration or missed opportunities and questions the commercial relevance of the learning horizon and synchronization assumptions. This paper therefore uses the environment as a controlled laboratory. In consumer markets, the agent may claim to serve the customer while vendor or provider incentives steer the choice. The action log shows prices or purchases. It does not show the lower, consumer-friendlier prices or options the agents considered and rejected. The important evidence sits inside the choice.

This paper asks what kind of record makes that evidence usable. It uses the Calvano environment because the welfare effect is measurable, the agents can learn harmful conduct without communication, and each decision can be recorded exactly. Consumer shopping agents and pricing algorithms are different, but both settings turn on the same question: what has to be recorded before an evaluator can tell whether an agent served the relevant user or was captured by another interest?

2. What the experiment shows

The experiment attaches decision evidence to each pricing decision. The outcome measure is the profit-normalized collusion index, denoted Δ . A value near 1 means monopoly-level profit, and a value near 0 means competitive-benchmark profit. Version 1.2 uses the published Calvano span-extension grid and makes each failed audit schedule exactly ten future restricted agent-decisions. Each reported cell contains 200 deterministic seeds. The expensive loop runs in Numba-compiled machine code; Python supplies the command-line, multiprocessing, and result-writing layers.

The design supports two empirical questions. First, does an outcome-contingent restriction change learned pricing relative to random interruption with matched realized restriction exposure? Second, does the rejected-alternatives trigger add an incremental effect when the outcome trigger and consequence remain fixed? The active outcome trigger reads rolling realized profit and consumer surplus. Those inputs can be computed from ordinary outcome logs, so its effect identifies outcome-contingent restriction rather than a decision-record effect.

The earlier “no-reader” label covered three distinct cells. The record-evidence-only cell computes the prediction and rolling evidence inputs and materializes complete sample records at session end, with zero forecast penalty, zero welfare weight, and zero audit probability. Two forecast-scoring cells reward more accurate probability forecasts. Accuracy is measured with the Brier score, the squared distance between the stated probability and what later occurs; lower scores are better. Each treatment uses the same seed as its control, so both sessions begin from the same random sequence. That paired comparison isolates whether activating this evidence-computation path changes conduct in this implementation; it does not test durable filing. The scoring comparisons test whether rewarding forecast accuracy changes calibration or conduct.

Table 1. Evidence production and changes to the learner’s objective.

Arm	Mean Δ	SE	Difference from its common-seed control	95% interval
Baseline	0.845	0.008	—	[0.829, 0.860]

Arm	Mean Δ	SE	Difference from its common-seed control	95% interval
Record-evidence only	0.845	0.008	0.000	[0.000, 0.000]
Forecast penalty 0.1	0.846	0.008	0.001	[-0.023, 0.025]
Forecast penalty 0.3	0.847	0.008	0.002	[-0.019, 0.024]
Consumer-surplus reward 0.1	0.597	0.015	-0.248	[-0.281, -0.215]
Welfare-in-reward 0.1	0.293	0.013	-0.552	[-0.583, -0.521]

The first result is exact. Under every common seed, activating the evidence-computation path leaves pricing unchanged. Forecast scoring also produces no detectable pricing change. The Brier score falls from early to late learning in every evidence-producing arm, including the arm with no scoring reward. The experiment therefore shows learning-period calibration; it does not show that scoring caused the improvement.

A record can become accurate while conduct remains unchanged

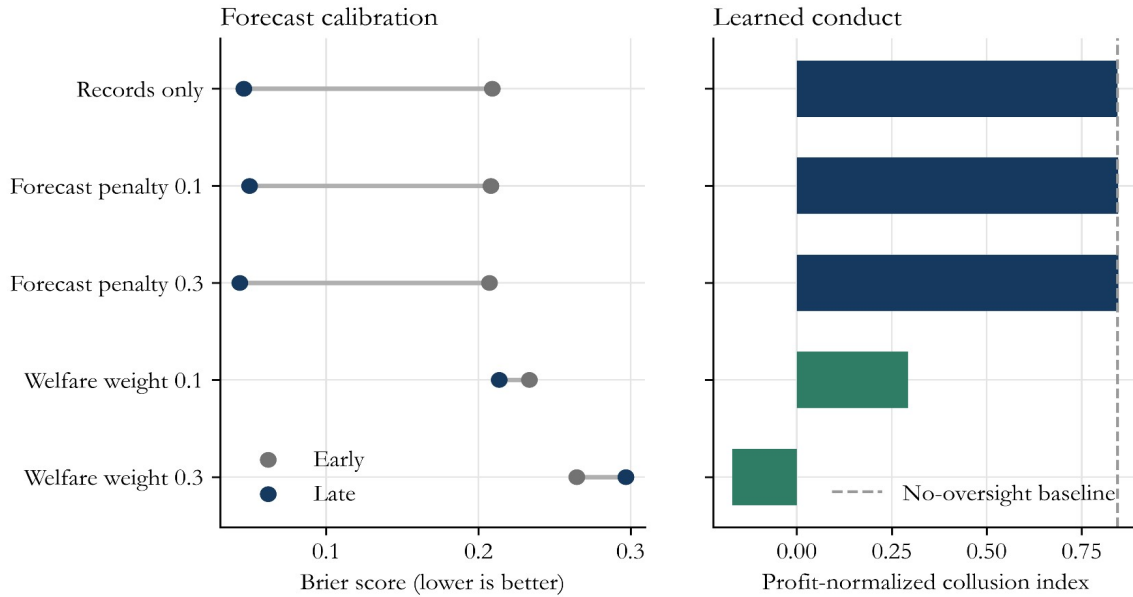


Figure 1. Forecast calibration improves during learning while evidence production and forecast scoring leave learned pricing unchanged. Welfare-weighted rewards change the private objective and the resulting conduct.

Changing the private objective has a large effect. The consumer-surplus and welfare-in-reward arms reduce Δ , with stronger weights pushing the index slightly below the competitive benchmark. These arms establish what happens when consumer welfare enters reward. They do not establish an effect of filing or reading a record.

Table 2. Outcome-contingent restriction and the matched random control.

Audit probability	Targeted Δ	Targeted exposure	Random Δ	Random exposure	Targeted minus random (95% interval)
0.02	0.793	0.098	0.802	0.098	-0.009 [-0.035, 0.018]

Audit probability	Targeted Δ	Targeted exposure	Random Δ	Random exposure	Targeted minus random (95% interval)
0.05	0.778	0.229	0.758	0.228	0.020 [-0.005, 0.045]
0.10	0.739	0.391	0.723	0.388	0.016 [-0.014, 0.046]
0.25	0.150	0.490	0.711	0.486	-0.561 [-0.603, -0.520]
0.50	0.036	0.433	0.733	0.436	-0.697 [-0.720, -0.674]

The exposure match is close at every audit probability. At the three lower probabilities, the intervals include zero. At the two highest probabilities, targeting produces much lower Δ than an equally frequent random interruption. Selection matters once the targeted rule is applied often enough in this environment. The response is nonmonotonic because restriction exposure peaks before the highest audit probability: more frequent audits can overlap existing restrictions and can alter the states that generate later failures.

Enforcement dose and learned pricing

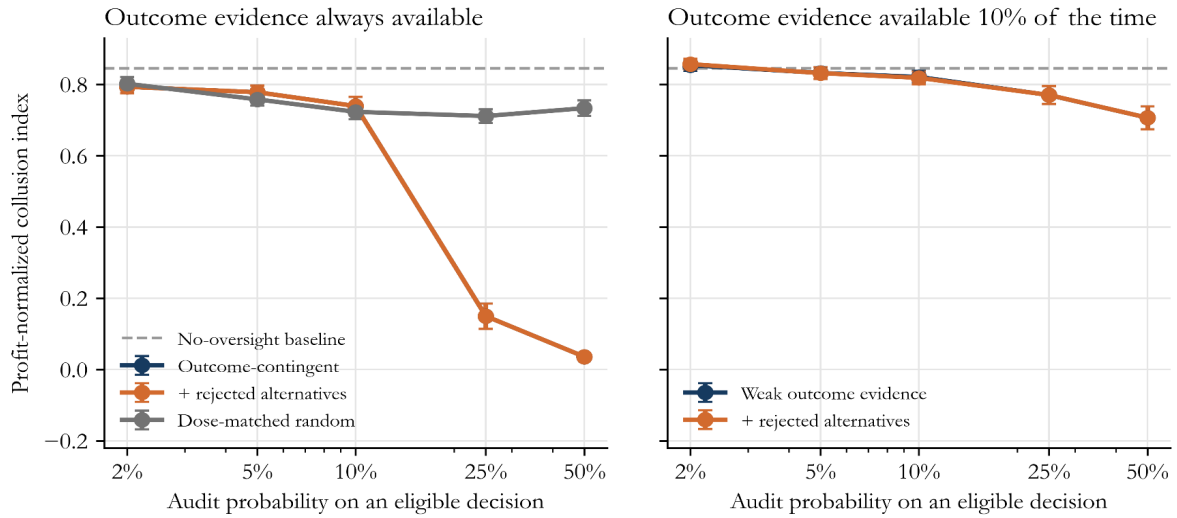


Figure 2. Outcome-contingent restrictions sharply reduce supracompetitive pricing at the two highest audit probabilities. The rejected-alternatives add-on overlaps the outcome-only path, while dose-matched random interruption does not reproduce the high-intensity effect.

Table 3. Tests of the record-specific and several-reader additions.

Comparison	Result	Mechanism evidence	Interpretation
Rejected alternatives added to the main outcome rule	Exact zero at every audit probability	Zero alternatives triggers	Implemented add-on is inactive
Rejected alternatives under weak outcome detection	No interval excludes zero	154 triggers across all cells; negligible exposure increment	No detectable incremental effect
Alternatives-only rule	Δ between 0.834 and 0.836	Exposure below 0.001%	Implemented trigger is too rare to change conduct
Several-reader composite added to the main rule	Exact zero at every audit probability	Zero composite triggers	Reader diversity is untested by this implementation

The rejected-alternatives result is a design failure with a useful diagnosis. A receipt omits the key counterfactual, but the implemented alternatives rule almost never finds a qualifying case. The result does not show that alternatives are unimportant. It shows that this threshold and persistence rule fails to create meaningful treatment exposure. Likewise, the several-reader composite never crosses its threshold. A test of reader diversity needs genuinely distinct rules that activate on some records.

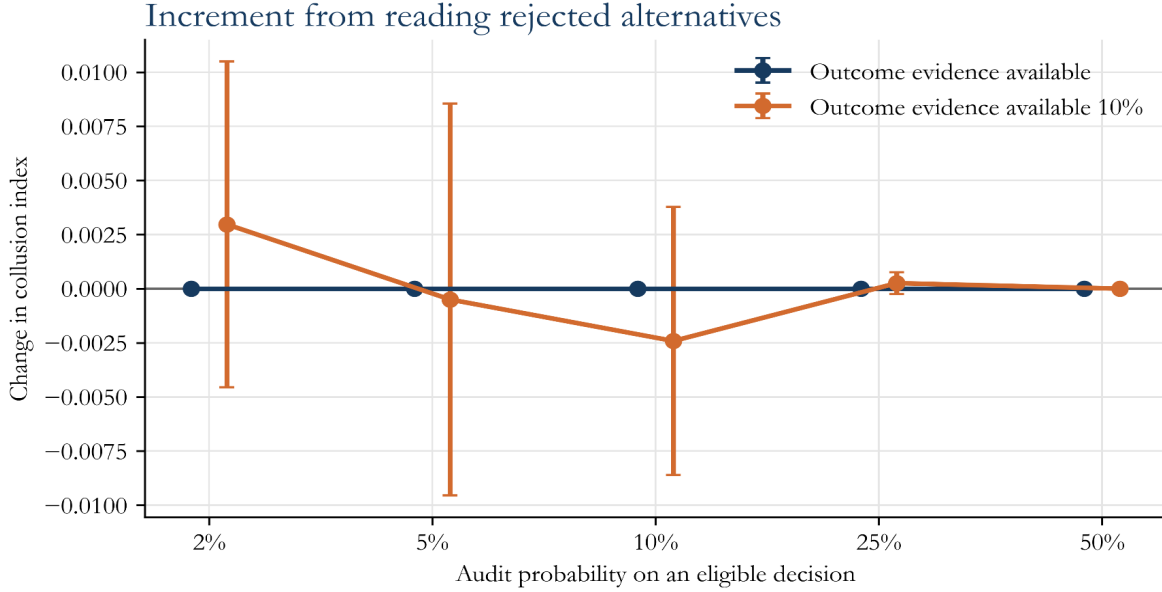


Figure 3. The incremental effect of the rejected-alternatives trigger is zero when outcome evidence is always available and statistically indistinguishable from zero under weak detection. Error bars are paired 95% intervals.

These findings establish the paper’s central distinction. Producing evidence changes no incentive. A reader or rule changes conduct only when it actually sees a qualifying event and can impose a consequence. The outcome rule succeeds at high intensity; the two record-specific additions fail to activate. DRCS-EC is therefore a proposed field specification and research program, not a validated consumer-agent intervention.

3. The DRCS-EC standard

The proposed DRCS-EC pilot would ask NIST and the EC to define a consumer-commerce decision-record profile for AI agents. It would require fields that matter for commerce: user preferences, sellers and products considered, total price, shipping terms, return terms, payment path, checkout friction, seller reliability, selected option, rejected options, and a declared system-welfare effect.

The translation from the pricing simulation is bounded. The laboratory uses seller-side pricing agents to isolate an oversight mechanism. DRCS-EC is a proposed consumer-commerce profile derived from that mechanism; field requirements, privacy design, verification, and interoperability require separate work.

Two EC fields matter most.

Rejected alternatives show whether the agent passed over a better deal. Without that field, an evaluator can only see the final purchase and guess. With that field, the evaluator can ask whether the agent ignored a better seller, lower total landed cost, better return policy, faster delivery option, safer payment path, or better task outcome.

The proposed system-welfare field would record effects beyond one checkout screen. In the Calvano study, the laboratory welfare quantity is consumer surplus. In commerce, candidate quantities include platform concentration, shipping waste from poorly coordinated delivery, return costs, privacy and security risk, and other effects that

matter beyond the single transaction. A pilot can standardize required inputs, field definitions, integrity rules, and conformance tests while comparing alternative measures.

In DRCS-EC, system welfare is the field that keeps the gatekeeper ratchet visible. It can record whether a choice concentrates purchases inside one marketplace, weakens multihoming, raises seller dependence, increases shipping waste, or routes transaction data into a gatekeeper's advantage. The standard should preserve common inputs and integrity rules so competing evaluators can compute their own rankings from the same evidence.

DRCS Core 1.0 is the broader Decision Accounting framework that EC adapts. Decision Accounting is a management science for material decisions across the economy: households, businesses, nonprofits, governments, and every organization that makes consequential choices. For the AI AGENT Act, EC is the operative standard. Core is the framework behind it. Both are published under the Apache 2.0 open-source license.

4. Why the standard is practical

The archived compute accounting measures aggregate simulation activity. The accelerated kernel retains numeric audit buffers and materializes complete record samples at session end. It does not construct, serialize, store, transmit, protect, or verify a complete record for every agent-decision. The manuscript therefore withdraws every per-record cost and record-count claim. A practical cost study must compare matched runs with no record construction, full in-memory construction, serialization, durable storage and indexing, integrity protection, and audit reading. It should report wall time, processor time, energy, bytes, and uncertainty across repeated runs.

The remaining question is institutional design. The experiment tests specified consequences in one laboratory. Congress and NIST can consider a common record profile and an evaluation program that compares reader rules, authority, privacy protections, and consequences. Reader diversity remains a proposed design feature whose behavioral effect requires an active treatment.

Potential market uses include platform evaluator tools, private ranking services, enterprise procurement screens, regulatory sampling, evidentiary review, and user assessment agents. A pilot should measure implementation cost, error rates, contestability, privacy risk, and behavioral effects for each use.

The SAPM framework motivates a contestability question. A single platform-run evaluator can become a gatekeeper control point. A pilot can test whether several readers using the same record improve detection or discipline when their rules and consequences actually differ.

The proposed public role is the record standard. Ranking, comparison, procurement screens, certification, enforcement theories, and evidentiary uses can be tested on top of that record. The record makes the decision inspectable. Whether heterogeneous readers improve loyalty incentives is an open empirical question.

5. Limits and conclusion

The experiment is deliberately narrow. It uses two Q-learning pricing agents in a corrected Calvano environment. The records are truthful by construction. The welfare calculation is available because consumer surplus is specified by the laboratory. Real consumer-agent markets will require tamper resistance, record verification, privacy controls, and domain-specific system-welfare definitions.

Those limits define the next engineering and standards tasks. The experiment tests outcome-contingent restriction against dose-matched interruption and tests the incremental rejected-alternatives trigger. It does not isolate durable filing or reader diversity. The consumer-market translation is a proposed DRCS-EC profile that preserves alternatives and declared system-welfare data for later evaluation.

The policy conclusion is conditional. Outcome-only ratings cannot reveal which alternatives an agent considered or rejected. DRCS-EC proposes fields for that evidence. The current design identifies a targeting effect at the two

highest audit intensities and no targeting effect at the three lower intensities. It also identifies a rejected-alternatives null caused by negligible trigger exposure. The implemented several-reader score adds no consequence, so reader diversity remains a Prat-motivated design proposal. NIST and the European Commission can use a pilot to test the record profile, active reader rules, verification, privacy, interoperability, consequences, and production cost before considering broader adoption.

Appendix A. The environment, the learner, and the audit rules

The appendices define the audit trail. Appendix A fixes the environment, learner, treatments, labels, and estimands. Appendix B reports robustness and mechanism checks. Appendix C specifies the replication and compute protocols. Appendix D documents the record schema.

This appendix collects the mechanical specification used in the main run. The purpose is auditability: a reader should be able to map every result table back to a parameter, formula, or trigger.

Table A1. Core simulation parameters.

Quantity	Value
Firms	$n = 2$
Marginal cost	$c = 1$
Product quality	$a = 2$
Outside good	$a_0 = 0$
Differentiation	$\mu = 0.25$
Discount factor	$\delta = 0.95$
Price grid	15 points on $[p^N - 0.1(p^M - p^N), p^M + 0.1(p^M - p^N)]$
State	Previous-period price pair
Q initialization	Discounted expected payoff against a uniform-random rival
Exploration	$\epsilon_t = \exp(-4 \times 10^{-6} t)$
Learning rate	$\alpha = 0.15$
Convergence rule	100,000 consecutive periods with unchanged greedy policy
Period cap	1,000,000,000 periods for C0, C1, C2, and C2W; 10,000,000 for audit and matched-interruption arms
Sessions per main cell	200 deterministic seeds

The benchmark values are $p^N = 1.4729$, $p^M = 1.9250$, $\pi^N = 0.2229$, $\pi^M = 0.3375$, $CS^N = 0.7151$, and $CS^M = 0.3271$. Displayed benchmark arithmetic uses these rounded four-decimal values; the code computes from unrounded solutions. The profit scale is $\pi^M - \pi^N = 0.1146$. The consumer-surplus gap is $CS^N - CS^M = 0.3880$.

To put each period's consumer surplus on a common scale, the model measures how far the chosen price is from the competitive benchmark, relative to the full competitive-to-monopoly consumer-surplus gap. A value of 0 is the competitive benchmark and -1 is the monopoly benchmark; the number is an accounting input to this simulation's reward, not a general welfare measure. The welfare term in the harm-priced arm is

$$\Delta W = \frac{CS(p) - CS^N}{CS^N - CS^M}.$$

In the welfare-in-reward arm the reward is $r' = \pi - \lambda_b \cdot \pi_{scale} \cdot b_t + \lambda_w \cdot \Delta W$, where b_t is the period- t Brier score, $\pi_{scale} = \pi^M - \pi^N$, λ_b is the prediction-score penalty weight, and λ_w is the system-welfare weight.

This reward changes the learner's private objective only in the harm-priced treatment. It is available here because the laboratory specifies consumer surplus exactly; consumer-agent deployments would require a separately justified welfare measure, evidence source, and uncertainty treatment.

For each session, the analysis reads `sessions_detail[].pi_bar` and the cell-level benchmark values `benchmarks.piN` and `benchmarks.piM`; computes $\Delta_s = (\pi_s - \pi^N) / (\pi^M - \pi^N)$; and reports the all-session mean with the across-session sample standard error. Subgroup columns are labeled **converged sessions** and **sessions ending at the period cap**. Each subgroup standard error uses the sample convention with denominator $n - 1$. Difference intervals use common-seed paired estimation. Forecast-calibration measures are reported separately for each evidence-producing arm.

Audit logic is deterministic conditional on the audit draw:

1. Draw an audit with probability q after the rolling window is populated.
2. The outcome-and-pattern audit triggers a failed audit if the window shows private profit above the competitive benchmark while consumer surplus is below the competitive benchmark, or if the record-pattern index G_i exceeds 2.0.
3. The alternatives-reading audit adds the alternatives component. A rejected-alternatives trigger requires a rejected action within 5% of the chosen action's Q-value range, a consumer-surplus improvement of at least 10% of the competitive-to-monopoly consumer-surplus gap, and the condition in more than 90% of the audit window.
4. The alternatives-only audit uses the same alternatives trigger and same consequence as the alternatives-reading audit, but disables the outcome and record-pattern triggers.
5. The several-reader audit adds the several-reader score. The score averages operator, truth, outcome, auditor, successor-maintainer, and benchmark-runner views. It triggers below 0.55.
6. Every trigger imposes the same consequence: after the trigger decision, the affected agent charges the competitive benchmark price $p^N = 1.4729$ for exactly ten future agent-decisions. A deterministic unit test must count all ten before any result is accepted.

The pattern check is intended to flag a pricing record that looks unusually coordinated or favorable to the seller across several clues at once. It combines those clues into one screening score so an audit can use a reproducible threshold; it does not infer collusion from any single clue.

The record-pattern index uses the trailing 100-period audit window. For each agent, the five raw components are mean predicted probability, action-sequence regularity, prediction negentropy, mean profit favorability $(\pi - \pi^N) / (\pi^M - \pi^N)$, and mean price-level favorability $(p - p^N) / (p^M - p^N)$. Each component is z-scored

across the two agents in the same window. The index is $G_i = \sqrt{\frac{1}{5} \sum_{k=1}^5 z_{ki}^2}$, and the outcome-and-pattern audit flags

agent i when $G_i > 2.0$. With two agents, each component's z-score has magnitude at most 1 when defined, so G_i cannot exceed 1. The stated threshold cannot activate in this implementation. The outcome-and-pattern arm therefore identifies the outcome trigger, and any record-pattern treatment requires a redesigned statistic or more agents.

The several-reader check asks whether one shared record would become more consequential when different legitimate readers assess different features of it. The composite below gives each of the six laboratory views equal weight solely to create a transparent test trigger; it is not a universal ranking of agents or readers. The several-reader score in the several-reader audit is

$$S = \frac{1}{6} (S_{operator} + S_{truth} + S_{outcome} + S_{auditor} + S_{successor} + S_{benchmark}).$$

A low score schedules the same ten-future-decision consequence used elsewhere in this simulation. Because the components and threshold are implementation choices, an active treatment would test only this composite.

The implemented several-reader score averages the last 15 rolling evidence states in the audit window. For state r and agent i , $S_{operator} = \min(1, \max(0, \pi_{ir} / \pi^M))$, $S_{truth} = \max(0, 1 - Brier_{ir} / 0.25)$, and $S_{outcome} = 1$ when $\Delta W_r \geq -0.02$. If $\Delta W_r < -0.02$, then $S_{outcome} = \max(0, 1 - (Q_{ir}^{max} - Q_{ir}^{min}) / (\pi^M - \pi^N))$. Auditor, successor-maintainer, and benchmark-runner scores equal 1.0 because the required machine-state inputs are present. The composite never crosses its trigger threshold and therefore adds no consequence. An experiment about reader diversity requires distinct active reader rules that can change treatment exposure.

The rerun uses the archive's arm labels. C0 is the no-oversight baseline; C1 puts consumer surplus directly into reward; C2 with zero forecast penalty is the record-evidence-only arm; the other C2 cells add forecast scoring; C2W adds system welfare to reward; C3a is the outcome-and-pattern audit, whose pattern component cannot activate with the present two-agent statistic; C3b adds the rejected-alternatives trigger; C3c uses the alternatives trigger alone; C3aP and C3bP are weak-outcome-detection variants; C3m adds the archived several-reader composite; and C3r is the dose-matched random-interruption control.

Table A2. Corrected-grid result-table specification.

Field	Required label and calculation
Headline estimand	All-session mean Δ
Uncertainty	Across-session sample standard error, denominator $n - 1$
Convergence subgroup	Converged sessions: mean, sample SE, and n
Period-cap subgroup	Sessions ending at the period cap: mean, sample SE, and n
Event count	Failed-audit episodes by trigger
Consequence count	Forced agent-decisions, measured from the realized restriction counter and adjusted for session-end truncation
Exposure	Forced agent-decisions divided by all agent-decisions
Status	Completed corrected-grid archive

Appendix B. Robustness and mechanism checks

Every check in this appendix uses the corrected price grid and an exactly ten-future-decision restriction. The all-session mean is the headline estimand. Converged-session and period-cap-session estimates are descriptive decompositions because convergence status is post-treatment.

Weak outcome detection

The weak-detection arms retain the same design while allowing the outcome component to verify harm with probability 0.10. Because the current two-agent record-pattern statistic cannot cross its threshold, C3aP remains a

weak outcome-trigger arm. C3bP adds the rejected-alternatives trigger. The paired C3bP-minus-C3aP intervals include zero at every audit probability. The weak outcome rule nevertheless shows a dose response: mean Δ falls from 0.854 at an audit probability of 0.02 to 0.706 at 0.50 as realized restriction exposure rises from 0.018 to 0.315.

Trigger attribution and exposure

The mechanism tables report these distinct objects:

- **Audit draws:** opportunities for review.
- **Failed-audit episodes:** trigger events, separated into outcome, record-pattern, rejected-alternatives, and several-reader channels.
- **Forced agent-decisions:** realized future decisions at the competitive benchmark, counted directly from the restriction state and adjusted for overlap and session-end truncation.
- **Exposure:** forced agent-decisions divided by total agent-decisions.

An episode count is never labeled as a period count. Scheduled exposure is used only for design calibration; realized exposure is used for the dose-matched comparison.

Rejected-alternatives threshold sensitivity

The threshold analysis preserves three prespecified variants:

Variant	Value tolerance	Consumer-surplus threshold
Loose value	0.100	0.100
Loose welfare	0.050	0.050
Strict both	0.025	0.200

The table reports the common-seed incremental effect of adding the alternatives trigger to the weak outcome rule, and the effect of the alternatives-only arm relative to baseline.

Audit probability	Variant	Alternatives added to weak outcome rule: Δ difference [95% CI]	Alternatives only: Δ difference [95% CI]
0.25	Loose value	+0.0038 [-0.0025, +0.0100]	-0.0077 [-0.0311, +0.0156]
0.25	Loose welfare	+0.0003 [-0.0002, +0.0008]	-0.0054 [-0.0281, +0.0172]
0.25	Strict both	+0.0000 [+0.0000, +0.0000]	-0.0070 [-0.0300, +0.0160]
0.50	Loose value	+0.0000 [+0.0000, +0.0000]	+0.0020 [-0.0209, +0.0249]
0.50	Loose welfare	+0.0166 [-0.0002, +0.0334]	-0.0048 [-0.0271, +0.0176]
0.50	Strict both	+0.0000 [+0.0000, +0.0000]	-0.0085 [-0.0315, +0.0145]

Every interval includes zero. The strict trigger produces no events in the weak-outcome arm at either audit rate. The loose-welfare variant at 0.50 produces 2,217 alternatives events in 12 of 200 sessions, but its incremental estimate remains imprecise and points toward more collusive pricing rather than less. In the alternatives-only cells, looser thresholds produce 510 to 3,364 events, yet realized restriction exposure remains between 0.000677% and 0.004460% of agent-decisions. The result is consistent across the threshold family: this alternatives rule does not generate enough consequential exposure to establish a pricing effect.

Dose-matched random interruption

C3r uses the same audit-rate grid and exactly ten future restricted decisions per event. Its event probability is calibrated so realized forced-agent-decision exposure matches C3a. The absolute exposure mismatch ranges from 0.0001 to 0.0037. The primary targeting contrast is C3a minus C3r at each audit rate. Its interval includes zero at 0.02, 0.05, and 0.10 and excludes zero by a wide margin at 0.25 and 0.50.

Learning parameters

The learning-parameter check reruns the baseline, record-evidence-only and forecast-scoring arms, C3a, and C3b at two alternative learning-rate and exploration-decay pairs.

Learning rate; exploration decay	Baseline Δ	Evidence-only Δ	Forecast-scored Δ	Outcome rule Δ	Outcome + alternatives Δ
0.10; 2×10^{-6}	0.868	0.868	0.851	0.040	0.040
0.20; 8×10^{-6}	0.831	0.831	0.836	0.034	0.034

The exact baseline/evidence-only equality survives both settings. The outcome rule remains strong at the highest audit probability, and the alternatives add-on again produces an exact zero increment because it does not trigger. Relative to the common-seed evidence-only control, forecast scoring changes Δ by -0.0165 [95% CI -0.0377, +0.0046] in the first setting and +0.0058 [-0.0161, +0.0277] in the second. The sign changes and both intervals include zero.

Several-reader composite

The C3m composite never activates and therefore adds no consequence beyond C3b. The corrected rerun is an exact diagnostic parity check: every C3m session matches C3b under its common seed. It cannot establish a reader-diversity effect. A future reader-diversity experiment must specify genuinely distinct active rules, verify that each rule changes exposure for some records, and compare the resulting consequence paths.

Required result-table labels

Earlier label	v1.2 label
Capped session	Session ending at the period cap
End-of-run Δ	State-at-stop Δ , with stop defined as convergence or period cap
Forced competitive price period	Failed-audit episode, when counting triggers
Forced competitive price period	Forced agent-decision, when counting restricted choices
SE	Across-session sample SE, unless another estimator is named

The version-matched results archive supplies every robustness estimate, interval, event count, exposure rate, convergence count, and figure reported here. ## Appendix C. Replication and compute protocol

The v1.1 result archive is superseded. The version-matched v1.2 package contains the regenerated cells, correction tests, tables, figures, and documented commands.

Correction gates

The rerun must satisfy these deterministic gates:

1. The 15 price points equal the Calvano span-extension grid from $p^N - 0.1(p^M - p^N)$ through $p^M + 0.1(p^M - p^N)$.
2. A single failed audit produces exactly ten future forced agent-decisions when the session continues long enough.
3. The duration test passes for C3a, C3b, C3m, C3c, and C3r.
4. Non-audit arms use the one-billion-period cap; audit and matched-interruption arms use the ten-million-period cap.
5. Every reported SE uses the across-session sample convention with denominator $n - 1$, unless the table names another estimator.
6. Trigger events, restricted decisions, scheduled exposure, and realized exposure remain separate variables.
7. C3r matches realized forced-agent-decision exposure within the prespecified tolerance.
8. Every figure and table reads only version-matched corrected result files.
9. The package imports from one internal module layout and every documented command exits successfully in a fresh environment.
10. The manifest, manuscript, result archive, figures, and checksums carry the same version.
11. Every production result cell reports **engine: numba**, and the packaged engine audit exits successfully; production runners stop if the compiled engine is unavailable.

Cell inventory

The complete rerun covers C0, the record-evidence-only C2 cell, forecast-scoring C2 cells, C1 and C2W reward changes, the full C3 audit-rate grid, weak-detection cells, rejected-alternatives-only cells, the several-reader parity check, dose-matched random interruption and its calibration, learning-parameter sensitivity, and rejected-alternatives threshold sensitivity. The record-path estimand is C2 with $\lambda_b = 0$, $\lambda_w = 0$, and $q = 0$ minus C0 under common seeds.

Cell counts, session counts, simulated-period counts, convergence counts, and treatment estimates are computed from the corrected archive. No v1.1 treatment estimate is carried forward.

Compute and record-production study

The accelerated archived kernel retained numeric rolling buffers and materialized complete record samples at session end. Conceptual agent-decisions cannot serve as a denominator for complete records that were never constructed and serialized. Aggregate simulation electricity cannot identify the incremental cost of the record layer.

A valid cost study will use matched runs and separately measure:

1. no record construction;
2. full 17-field in-memory construction;
3. serialization;
4. durable storage and indexing;
5. integrity protection or signing, if claimed;
6. record retrieval and audit reading.

For each stage, the package will report repeated-run wall time, processor time, energy, bytes written, record count, hardware and software environment, electricity-price basis, and uncertainty. Per-record cost and scalability claims remain withdrawn until that study exists.

Appendix D. Decision-record schema

Table D1 lists the full record schema represented by the simulation and materialized in its sample records; Section 3 develops the consumer-agent analogue.

Table D1. Decision-record components.

Component	What it answers in this experiment
Who	Which firm-agent made the pricing decision
What	The price the agent chose
When	Session, period, and treatment arm
Where	The Calvano duopoly environment and state
Why	The objective, reward signal, and value assigned to the chosen price
Evidence	The observed state and value table available at the moment of choice
Authority	The policy authorized to choose, and any forced competitive price constraint
Training	Learning rate, exploration schedule, and training period
Review	Scheduled review path and forced review conditions
Stakeholders	Firms, consumers, auditors, future readers, and the market
Consequences	Private profit and consumer-surplus consequence
Constraints	Price grid, exploration rule, forced competitive price rule, and convergence rule
Uncertainty	Exploration state and value dispersion
Communication	Which readers can receive or inspect the record
Alternatives	Rejected prices, their value to the agent, and their effect on consumers
Prediction	A scoreable forecast later evaluated with the Brier score
System welfare	The direction and magnitude of market-level harm or benefit

In the emitted JSON, the alternatives component stores `chosen_action` as an integer, `chosen_q` as a floating-point Q-value, `alternatives_summary.serious_options` as a list of rejected action objects, and the implementation key `alternatives_ledger.top_rejected` as a list with `action_index`, `price`, `q_value`, `delta_q_from_chosen`, `near_equivalent`, and `cs_delta_vs_chosen`. The prediction component stores a scoreable forecast: statement text, confidence percentage, binary observed outcome, and Brier score. The system-welfare component stores the system identifier, direction, magnitude on the competitive-to-monopoly consumer-surplus gap, monitoring signals, and reversal trigger. The version-matched package will include corrected-grid sample records tied to the rerun manifest.

Data and replication availability

The v1.1 replication package is superseded. The version-matched v1.2 manuscript, code, result archive, figures, tables, environment specification, manifest, and checksums are available from the project page: <https://decisionaccounting.org/calvano-dracs/>. The direct archive is <https://decisionaccounting.org/research/calvano-dracs/calvano-dracs-v1.2-replication-package.zip>. The package gives stable paths for the paper source and the main corrected results. No one-off request to the author is required.

Statements and declarations

Funding. No external funding supported this research.

Competing interests. The author founded the Center for Decision Accounting and developed the Decision Record & Control Standard evaluated in this paper. He may receive professional, reputational, or commercial benefits if Decision Accounting or related implementations are adopted. The Center is independent and has no affiliation with or endorsement from NIST, the European Commission, Microsoft, Anthropic, Google, OpenAI, or the authors of the Calvano environment.

Generative-AI disclosure. During manuscript preparation, the author used Anthropic Claude, OpenAI ChatGPT and Codex, and Google Gemini and Deep Research to assist with literature discovery, code review, consistency checking, adversarial review, and editorial revision. The author selected the research question and experimental design, reviewed the cited sources and computational outputs, revised the manuscript, and takes responsibility for the content.

References

- Abada, I., and X. Lambin. 2023. “Artificial Intelligence: Can Seemingly Collusive Outcomes Be Avoided?” *Management Science* 69 (9): 5042-5065.
- Agentic Commerce Protocol. 2026. Open specification, Apache 2.0 license, maintained by OpenAI and Stripe. <https://github.com/agentic-commerce-protocol/agentic-commerce-protocol>. Accessed July 8, 2026.
- Aguirre, A., G. Dempsey, H. Surden, and P. B. Reiner. 2020. “AI Loyalty: A New Paradigm for Aligning Stakeholder Interests.” *IEEE Transactions on Technology and Society* 1 (3): 128-137. <https://doi.org/10.1109/TTS.2020.3013490>.
- Amazon.com, Inc. v. Perplexity AI, Inc. 2025. Complaint filed November 4, 2025, U.S. District Court for the Northern District of California.
- Asker, J., C. Fershtman, and A. Pakes. 2022. “Artificial Intelligence, Algorithm Design, and Pricing.” *AEA Papers and Proceedings* 112: 452-456.
- Asker, J., C. Fershtman, and A. Pakes. 2024. “The Impact of Artificial Intelligence Design on Pricing.” *Journal of Economics & Management Strategy* 33 (2): 276-304. <https://doi.org/10.1111/jems.12516>.
- Assad, S., R. Clark, D. Ershov, and L. Xu. 2024. “Algorithmic Pricing and Competition: Empirical Evidence from the German Retail Gasoline Market.” *Journal of Political Economy* 132 (3): 723-771.
- Banchio, M., and A. Skrzypacz. 2022. “Artificial Intelligence and Auction Design.” In *Proceedings of the 23rd ACM Conference on Economics and Computation*, 30-31. New York: Association for Computing Machinery.
- Becker, G. S. 1968. “Crime and Punishment: An Economic Approach.” *Journal of Political Economy* 76 (2): 169-217.
- Brown, Z. Y., and A. MacKay. 2023. “Competition in Pricing Algorithms.” *American Economic Journal: Microeconomics* 15 (2): 109-156.
- Calvano, E., G. Calzolari, V. Denicolò, and S. Pastorello. 2020. “Artificial Intelligence, Algorithmic Pricing, and Collusion.” *American Economic Review* 110 (10): 3267-3297.
- Calvano, E., G. Calzolari, V. Denicolò, and S. Pastorello. 2021. “Algorithmic Collusion with Imperfect Monitoring.” *International Journal of Industrial Organization* 79: 102712. <https://doi.org/10.1016/j.ijindorg.2021.102712>.
- Calvano, E., G. Calzolari, V. Denicolò, and S. Pastorello. 2023. “Algorithmic Collusion: Genuine or Spurious?” *International Journal of Industrial Organization* 90: 102973. <https://doi.org/10.1016/j.ijindorg.2023.102973>.
- Calvano, E., G. Calzolari, V. Denicolò, J. E. Harrington Jr., and S. Pastorello. 2020. “Protecting Consumers from Collusive Prices Due to AI.” *Science* 370 (6520): 1040-1042.
- CNBC. 2026. “Amazon Wins Court Order to Block Perplexity’s AI Shopping Agent.” March 10, 2026. <https://www.cnbc.com/2026/03/10/amazon-wins-court-order-to-block-perplexitys-ai-shopping-agent.html>. Accessed July 8, 2026.

- Chan, A., C. Ezell, M. Kaufmann, K. Wei, L. Hammond, H. Bradley, E. Bluemke, N. Rajkumar, D. Krueger, N. Kolt, L. Heim, and M. Anderljung. 2024. "Visibility into AI Agents." In *Proceedings of the 2024 ACM Conference on Fairness, Accountability, and Transparency*, 958-973. New York: Association for Computing Machinery. <https://doi.org/10.1145/3630106.3658948>.
- den Boer, A. V., J. M. Meylahn, and M. P. Schinkel. 2026. "Artificial Collusion: Examining Supracompetitive Pricing by Q-Learning Algorithms." *Management Science*, published online June 9. <https://doi.org/10.1287/mnsc.2024.08557>.
- Ezrachi, A., and M. E. Stucke. 2016. *Virtual Competition: The Promise and Perils of the Algorithm-Driven Economy*. Cambridge, MA: Harvard University Press.
- European Union. 2024. Regulation (EU) 2024/1689 (Artificial Intelligence Act), Article 12. *Official Journal of the European Union*, July 12. <https://eur-lex.europa.eu/eli/reg/2024/1689/oj>.
- Federal Trade Commission and U.S. Department of Justice. 2026. "Federal Trade Commission and Department of Justice Seek Public Comment for Guidance on Business Collaborations." Press release, February 2026. <https://www.ftc.gov/news-events/news/press-releases/2026/02/federal-trade-commission-department-justice-seek-public-comment-guidance-business-collaborations>. Accessed July 7, 2026.
- Google. 2026. "Product Data Specification" and "Set Up Your Return Policies for Shopping Ads and Free Listings." Google Merchant Center Help. <https://support.google.com/merchants/answer/7052112>. Accessed July 8, 2026.
- Google Cloud Marketplace. 2026. "Offer AI Agents through Google Cloud Marketplace." <https://docs.cloud.google.com/marketplace/docs/partners/ai-agents>. Accessed July 8, 2026.
- Hansen, K. T., K. Misra, and M. M. Pai. 2021. "Frontiers: Algorithmic Collusion: Supra-competitive Prices via Independent Algorithms." *Marketing Science* 40 (1): 1-12.
- Harrington, J. E., Jr. 2018. "Developing Competition Law for Collusion by Autonomous Artificial Agents." *Journal of Competition Law & Economics* 14 (3): 331-363.
- Harrington, J. E., Jr., and D. Imhof. 2022. "Cartel Screening and Machine Learning." *Stanford Computational Antitrust* 2: 133-154. <https://law.stanford.edu/wp-content/uploads/2022/08/harrington-imhof-2022.pdf>.
- Hartline, J. D., S. Long, and C. Zhang. 2024. "Regulation of Algorithmic Collusion." In *Proceedings of the 2024 Symposium on Computer Science and Law*, 98-108. New York: Association for Computing Machinery. <https://doi.org/10.1145/3614407.3643706>.
- Hartline, J. D., C. Wang, and C. Zhang. 2025. "Regulation of Algorithmic Collusion, Refined: Testing Pessimistic Calibrated Regret." In *Proceedings of the 2025 Symposium on Computer Science and Law*. New York: Association for Computing Machinery. <https://doi.org/10.1145/3709025.3712217>.
- Hayek, F. A. 1945. "The Use of Knowledge in Society." *American Economic Review* 35 (4): 519-530.
- Johnson, J. P., A. Rhodes, and M. Wildenbeest. 2023. "Platform Design When Sellers Use Pricing Algorithms." *Econometrica* 91 (5): 1841-1879.
- Klein, T. 2021. "Autonomous Algorithmic Collusion: Q-Learning under Sequential Pricing." *RAND Journal of Economics* 52 (3): 538-558.
- Kolt, N. Forthcoming. "Governing AI Agents." *Notre Dame Law Review* 101. <https://arxiv.org/abs/2501.07913>.
- Miklos-Thal, J., and C. Tucker. 2019. "Collusion by Algorithm: Does Better Demand Prediction Facilitate Coordination Between Sellers?" *Management Science* 65 (4): 1552-1561.
- Microsoft. 2026. "Microsoft Marketplace: Cloud Solutions, AI Apps, and Agents." <https://www.microsoft.com/en-us/marketplace>. Accessed July 8, 2026.
- National Institute of Standards and Technology. 2026. "AI Agent Standards Initiative." <https://www.nist.gov/artificial-intelligence/ai-agent-standards-initiative>. Accessed July 7, 2026.
- OECD. 2017. *Algorithms and Collusion: Competition Policy in the Digital Age*. Paris: OECD Publishing.
- OpenAI. 2025. "Buy It in ChatGPT: Instant Checkout and the Agentic Commerce Protocol." <https://openai.com/index/buy-it-in-chatgpt/>. Accessed July 8, 2026.
- Prat, A. 2005. "The Wrong Kind of Transparency." *American Economic Review* 95 (3): 862-877.

- Postnieks, E. 2026a. “The Missing System Theorem.” Working paper.
- Postnieks, E. 2026b. “Decision Accounting: From Nudge to Guardrail.” Working paper.
- Postnieks, E. 2026c. “Multi-Audience Transparency and the Dissolution of Conformism: General Equilibrium Foundations for Decision Accounting.” Working paper.
- Postnieks, E. 2026d. “Sixty Years of Independent Convergence: Fourteen Base Documentation Fields Across a Twelve-Regime Core and Four-Regime Extension Corpus.” Working paper.
- Postnieks, E. 2026e. “From Nudge to Guardrail: Behavioral Foundations of Mandatory Pre-Decisional Documentation.” Working paper.
- Postnieks, E. 2026f. “Decision Records as Multi-Principal Governance Mechanisms: The Decision Record & Control Standard v1.0.” Working paper.
- Postnieks, E. 2026g. “Applying the System Asset Pricing Model to Platform Monopoly: Measuring the System Welfare Cost of Big Tech Acquisitions and Gatekeeper Rent Extraction.” Working paper.
- Schema.org. 2026. “MerchantReturnPolicy.” <https://schema.org/MerchantReturnPolicy>. Accessed July 8, 2026.
- Salesforce. 2025. “Salesforce Unveils AgentExchange, the Trusted Marketplace and Community for Agentforce.” Press release, March 4. <https://www.salesforce.com/news/press-releases/2025/03/04/agentexchange-announcement/>. Accessed July 8, 2026.
- Schildknecht, J. 2026. “Tacit Algorithmic Collusion: A Replication Study of Calvano et al. (American Economic Review, 2020).” *Journal of Comments and Replications in Economics* 5 (7). <https://doi.org/10.18718/81781.58>.
- Stripe. 2025. “Stripe Powers Instant Checkout in ChatGPT and Releases Agentic Commerce Protocol Codeveloped with OpenAI.” <https://stripe.com/newsroom/news/stripe-openai-instant-checkout>. Accessed July 8, 2026.
- U.S. Department of Energy, Energy Information Administration. 2026. “Electric Power Monthly, Table 5.6.A.” https://www.eia.gov/electricity/monthly/epm_table_grapher.php?t=epmt_5_6_a. Accessed July 7, 2026.
- U.S. Department of Justice. 2025. “Justice Department Requires RealPage to End the Sharing of Competitively Sensitive Information and Alignment of Pricing Among Competitors.” Press release, November 24. <https://www.justice.gov/opa/pr/justice-department-requires-realpage-end-sharing-competitively-sensitive-information-and>. Accessed July 7, 2026.
- Warner, M. R. 2026. “Warner Unveils Discussion Draft of Legislation to Create Innovative Market for Secure Artificial Intelligence Agents.” Press release, June 29. <https://www.warner.senate.gov/newsroom/press-releases/warner-unveils-discussion-draft-of-legislation-to-create-innovative-market-for-secure-artificial-intelligence-agents/>. Accessed July 8, 2026.

Acknowledgments

I am responsible for the design, simulations, interpretation, policy translation, replication package, and any remaining errors.

About the author

Erik Postnieks is the founder of the Center for Decision Accounting, an independent research center developing the Decision Accounting, Missing System Theorem, and System Asset Pricing Model research programs. He holds a B.A. in Economics from Cornell University and an M.B.A. from the Amos Tuck School of Business at Dartmouth College. His work focuses on governance architecture, system-welfare measurement, decision records, and the institutional conditions under which private incentives fail to price system consequences.