

Refundable Stake Does Not Expand Deterrent Capacity in Long-Lived Delegation: Collateral, Premia, and Pseudonymous Markets

Wulf A. Kaal, Ph.D.*

August 2026

Abstract

*Professor of Law, University of St. Thomas School of Law (Minneapolis). *Nature and scope of claims.* This Article is a theory paper: a reduced-form incentive analysis of refundable collateral in long-lived delegated relationships among pseudonymous parties. It distinguishes three classes of statement and labels each in place: the author's results whose algebra survived the adversarial audit protocol described below (the capacity frontier, its cell slopes, and the admission conditions); conditional results whose validity depends on stated open modeling choices (the market layer and any equilibrium-selection role for entry fees); and conjectures (the behavioral decomposition for observed misconduct rates). Complete formal proofs are not contained in this draft. After the author's initial formulation, the claims below were stress-tested in a nineteen-round adversarial protocol in which three frontier language models (Claude, Grok and Codex), under the author's direction, alternately attempted to identify counterexamples, algebraic errors, and overstatements. The surviving statements are those that neither system broke, and the audit trail is preserved in the author's research archive. That protocol is neither peer review nor a substitute for complete formal proofs. Every reference was verified against publisher or authoritative bibliographic records, with resolving links, on August 20, 2026, and each is cited only for the specific statement the verification confirmed it supports. Nothing in this Article predicts the performance of any commercial deployment or constitutes investment, legal, or technical advice. Field environments may diverge from the model in ways the analysis prices only partially, including adversarial populations, filing and adjudication quality, market tightness, real rather than modeled stakes, regulatory and jurisdictional constraints, and the operational, governance, and incentive choices of commercial principals that are not within the author's control. Any party considering reliance on this work for investment, deployment, or other non-academic purposes should conduct independent verification under its own conditions. *Conflict-of-interest disclosure.* The author is simultaneously the theorist, the protocol architect, and the empiricist for the research program of which this Article is part. The work was conducted on compute infrastructure owned by the author. No external funding supported this research. *Non-reliance.* The author is not making, and this paper does not constitute, any forward-looking statement, prediction of, or representation about the performance of any commercial instantiation, token offering, or investment vehicle. This paper is not

Practitioners of delegated pseudonymous markets, from proof-of-stake validation to agent economies, commonly assume that a slashable, refundable bond deters misconduct: the operator who cheats loses the stake, so a larger stake buys more honesty. This draft reports the central result of a formal project testing that assumption: in the strict long-match limit, increasing refundable escrow does not expand the credible capacity frontier. In a model where identities can be discarded and re-minted at negligible cost, a refundable bond released at exit pads the operator’s walk-away value by the full release-discounted principal; custody carry is an additional tax. In the strict long-match limit, refundable escrowed principal weakly contracts the credible capacity frontier in all four cells of a two-by-two regime map spanning exit versus in-place misconduct and leaky versus captive stake: the contraction is strict in both exit cells and strict in both stay cells if and only if the adjudication probability is positive. The relationship premium supplies the continuation-value component of deterrence. The capacity frontier is identically $\theta X \leq P + \varphi(d_h - R)F$ against exit misconduct and, at the baseline with no claw-back and no continuation damage, $\theta X \leq pP + p[d_h - (1 - s)\delta_E]F$ against in-place misconduct. Those displays are accounting identities at a given (P, F) ; the envelope derivatives in F are the slopes reported below, which already include the decline of P in F . Away from the limit, escrow has a positive marginal effect only under turnover-funded admission conditions derived in closed form; at the ideal leaky-exit corner this requires timely filing coverage above a floor that exceeds one in the strict long-match limit with positive carry. The result yields a regime-general capacity prediction; deriving predictions for observed misconduct rates requires an additional behavioral model, offered here as a conjecture. A second, conditional result endogenizes the market outside option and bounds admissible entry fees; whether a small fee also selects among formation equilibria remains open pending an extensive-form formation game.

Keywords: reputation, relational contracts, collateral, slashing, pseudonymity, mechanism design, delegated custody.

JEL codes: C73, D82, D86, L14, G23.

investment, legal, or technical advice. The author requests prior written review before any portion of this paper is quoted, paraphrased, or incorporated into offering materials, marketing materials, or other public commercial statements; this request does not purport to restrict quotation or other use permitted by applicable law. No advisor, co-founder, or promoter of any commercial party is authorized by the author to make representations sourced from this paper. This is a working paper. It has not been peer reviewed and remains subject to revision. Comments welcome: wulf@wulfkaal.com.

1 Introduction

The smarter agents have arrived, and the markets that employ them have converged on a single disciplinary device: the bonded stake. An operator posts collateral with a mechanism; misconduct, once adjudicated, burns it; honesty returns it. A common rationale for slashing protocols, and for informal arguments that pseudonymous agents will behave because they have skin in the game, is that posting more refundable stake increases deterrence.

The paper maintains that the rationale fails where it matters most. The flaw is not detection, latency, or adjudication error, although each is priced in the model. The flaw is structural: a bond that is refunded on clean exit is an asset the operator carries out the door. It raises the operator’s walk-away value by the honest-release value of the principal, while misconduct changes recovery only through the filing-, freeze-, adjudication-, and slash-weighted recovery wedge; custody carry further reduces capacity. Release pads the walk. For long-lived delegation with positive carry, the marginal capacity effect is strictly negative in both exit cells and, whenever the adjudication probability is positive, in both stay cells: refundability raises walk-away value and imposes carry, so additional principal fails to expand long-match capacity, and the relationship premium, the value of staying over leaving, is the scalable continuation-value component of deterrence.

This is a classification, not a lament. It is also the formalization of a position the author’s scholarship has maintained for eight years: that reputation, the standing value of a relationship, is the accountability substrate, and collateral is not (Calcaterra, Kaal, and Andrei 2018; Kaal 2019, 2021, 2024, 2026). Proposition 1 is this paper’s operational envelope for that position; it does not claim that the cited papers derived the inequality. Once the premium is identified as the scalable source of long-match deterrent capacity, the design question changes. The productive questions are no longer how large the bond must be, but how premia are manufactured, how the release clock can discriminate between clean and contested exits, which misconduct technologies leave the stake captive, and what entry pricing does to the market equilibrium that generates the premia. The paper answers the envelope questions in closed form; the fee-selection question is left open.

Three results organize the contribution. First, the dominance result (Theorem 1): in the long-match limit, refundable principal does not expand the capacity frontier in any regime cell, and strictly contracts it in the exit cells and in the stay cells with positive adjudication probability, with no escape hatch on that domain; the two apparent hatches, short cycles and claim-conditioned release freezes, open only under exact turnover-funded admission conditions that the long-match limit itself excludes (Corollary 1). Second, the

instrument results: a filing floor (Proposition 2) showing that claim-conditioned freezes require timely filing coverage exceeding a duration-determined threshold, so that unnoticed misconduct, which exits disguised as honesty, bounds what any release-clock policy can achieve; and a captive-stake variant showing that misconduct punishable in place eliminates the padding leak at the price of slicing the premium by the adjudication probability. Third, a conditional market-layer analysis (Proposition 3) yields a reduced form for the active-branch premium and an admissible fee interval; the fee does not enter the paper’s incentive constraint, while its possible equilibrium-selection role remains unresolved pending a specified extensive-form formation game.

2 Related literature

The paper stands on five strands and claims priority over none of them. Classical reputation theory (Kreps and Wilson 1982; Fudenberg and Levine 1989) sustains cooperation through incomplete information about types; the folk theorem tradition (Friedman 1971; Fudenberg and Maskin 1986) sustains it through continuation value. Both presuppose a persistent actor. Friedman and Resnick (2001) posed the cheap pseudonym problem and analyzed entry fees among possible remedies. Deterrence through expected sanction is Becker (1968); bonds and quality premia are Klein and Leffler (1981) and Shapiro (1983), while Telser (1980) supplies the self-enforcement mechanism through future-relationship value; the failure of deferred-wage bonding on timing grounds is Akerlof and Katz (1989), and the employer-side moral hazard objection to explicit worker bonds is pressed in Shapiro and Stiglitz (1985), with Carmichael (1985) as the pro-bonding counterargument. Relational enforcement under two-sided limited commitment is developed by MacLeod and Malcomson (1989) and Levin (2003), while Thomas and Worrall (1994) and Ray (2002) analyze backloading and the time structure of self-enforcing agreements. Cooperation under anonymous rematching is Ghosh and Ray (1996) and Kranton (1996a), and the erosion of relational exchange by attractive outside markets is Kranton (1996b); the market layer below reproduces that tightness effect with an exact rate. Collateral as a substitute for limited commitment is Hart and Moore (1994), Holmstrom and Tirole (1997), and Rampini and Viswanathan (2010, 2013); dynamic agency with cash state variables is DeMarzo and Fishman (2007) and Biais, Mariotti, Plantin, and Rochet (2007); the composition of formal and relational incentive instruments is Baker, Gibbons, and Murphy (1994). The institutional setting and the reputation-as-capital framing draw on the author’s prior work: Calcaterra, Kaal, and Andrei (2018) on infrastructure for

measuring domain-specific reputation; Kaal (2019) on why decentralized reputation verification systems are needed; Kaal (2021) on reputation as capital; Kaal (2024) on AI governance via Web3 reputation systems; and Kaal (2026) on the evolution of domain-specific reputation systems. The author’s proof-of-stake design line supplied a mechanism intuition, not a derivation of the present envelope: Calcaterra and Kaal (2018b) argue that stakes denominated in non-fungible reputation tokens incentivize long-term probity and eliminate short-term arbitrage opportunities that fungible cryptocurrency stakes permit; Kaal (2021b) proposes separating block consensus from rewards and making rewards a function of node reputation; and Kaal (2025) develops the foundations of that Secure Proof of Stake design. The present paper supplies the incentive accounting for why those design choices point in the right direction within its model. What the present envelope isolates is the padding netting: a refundable bond’s marginal contribution to the incentive constraint is on-path release value net of carry and thief recovery, which in the long-match limit is nonpositive in every cell and is positive away from the limit only under the turnover-funded admission conditions below; any claim of novelty for that accounting remains subject to the primary-source verification noted above.

3 Model

A delegated match pairs a capital-poor operator with a backing owner under a mechanism that custodies an operator-attributable escrow balance F . Time is continuous with discount rate r . The match ends by clean voluntary exit (hazard λ_e), by adjudicated failure (hazard λ_d , escrow burned), or by the operator’s misconduct.

Release clocks discriminate by claim status. An honest exit recovers d_h per escrowed dollar, where d_h mixes the clean release discount with false-freeze contamination and false adverse findings. An exiting thief recovers d_θ per dollar, which mixes three branches: unfiled misconduct (probability $1-\varphi$, where φ is timely pre-release filing coverage) recovers as an honest look-alike at d_h ; filed but unupheld claims recover at the frozen clock discount δ_f ; upheld claims forfeit the slashed fraction s and recover $(1-s)$ at the post-ejection discount δ_E . In symbols, with A the duration factor:

$$A = \frac{1}{r + \lambda_e + \lambda_d}, \quad C_F = r_c A, \quad q = \lambda_e d_h A, \quad (1)$$

$$d_\theta = (1 - \varphi) d_h + \varphi(1 - u) \delta_f + \varphi u(1 - s) \delta_E, \quad p = \varphi u, \quad (2)$$

where u is the uphold probability conditional on timely filing, r_c the custody carry rate,

q the on-path release value per escrowed dollar, and C_F the marginal carry cost. Let $R = (1 - u)\delta_f + u(1 - s)\delta_E$ denote the filed-thief recovery mixture, so $d_\theta = (1 - \varphi)d_h + \varphi R$ and $d_h - d_\theta = \varphi(d_h - R)$.

Deterrence runs through the operator's promised continuation. The credible promise is bounded by both parties' walk options. Each regime's incentive constraint compares appropriable exposure θX with the relationship premium P , defined as total match value net of both walk values, plus that regime's explicitly defined adjudication-weighted forfeiture value. The three governing envelope derivatives are:

$$\frac{\partial \text{cap}_x}{\partial F} = q - C_F - d_\theta \quad (\text{exit-theft slope}), \quad (3)$$

$$\frac{\partial \text{cap}_i}{\partial F} = p(q - C_F - (1 - s)\delta_E) \quad (\text{in-place slope}), \quad (4)$$

$$q - C_F - d_h = -(1 - \lambda_e A)d_h - C_F < 0 \quad (\text{participation slope}), \quad (5)$$

with the participation slope negative whenever recovery or carry is nondegenerate.

4 The premium principle and the main result

Proposition 1 (The premium principle). *Given $P = T - W_{\text{op}} - W_{\text{own}}$ and $W_{\text{op}} = d_h F + \Omega$, the exit-misconduct capacity frontier is identical to*

$$\theta X \leq P + (d_h - d_\theta)F, \quad d_h - d_\theta = \varphi(d_h - R), \quad (6)$$

and the in-place frontier, at $\chi = 0$ and $\Delta_U = 0$, is identical to

$$\theta X \leq pP + p[d_h - (1 - s)\delta_E]F. \quad (7)$$

These equalities substitute the walk floors into P ; they are not envelope derivatives in F .

In (6), P is the relationship premium and $(d_h - d_\theta)F$ the recovery wedge on stake; in (7), pP is the adjudication-weighted premium.

What keeps the operator honest is the value of the relationship itself: the premium P of staying over leaving, composed of future earnings, standing, and the cost of starting over. The stake F enters these fixed- (P, F) identities only through the recovery wedges. When F varies, the total derivatives of the capacity frontiers include the induced decline in P ; in the strict long-match limit those total derivatives are the nonpositive slopes reported

in Theorem 1. Thus the premium is the only component that expands the long-match frontier. Stake earns a positive marginal effect only under the cell-specific turnover-funded admission conditions of Corollary 1. In the leaky-exit cell that requirement is a conjunction: $d_h > R$ and $\varphi > \varphi^*$. The inequality $\varphi^* < 1$ (equivalently $q > C_F + R$, given $d_h > R$) is feasibility of the filing threshold, not satisfaction of it. Short cycles can separately meet the captive-exit and stay-cell conditions without any filing requirement. The paper's contribution is the exact pricing of those exceptions; its message is the primacy of the premium. This envelope is the present paper's counterpart of the author's reputation-as-capital framing (Calcaterra, Kaal, and Andrei 2018; Kaal 2021, 2024), not a result those papers stated.

Theorem 1 (Long-match principal dominance). *In the strict long-match limit ($q \rightarrow 0$) with positive carry, refundable escrowed principal weakly contracts the credible capacity frontier in all four cells of the regime map. The four cell slopes are*

$$\left(-C_F - d_\theta, \quad -C_F, \quad -p[C_F + (1-s)\delta_E], \quad -pC_F \right) \quad (8)$$

for the leaky-exit, captive-exit, leaky in-place, and captive in-place cells respectively: non-positive everywhere, strictly negative in both exit cells, strictly negative in both stay cells whenever $p > 0$, and flat in the stay cells at $p = 0$. No escape hatch exists on this domain: if $d_h \leq R$, with R the filed-thief recovery mixture defined above, the release-clock wedge is impossible outright; if $d_h > R$, the wedge's filing threshold exceeds one, because $\varphi^ > 1$ exactly when $q < C_F + R$, which positive carry guarantees in the limit. Short cycles are excluded by the limit itself.*

The economic content is the padding netting. In the leaky-exit cell, an escrowed dollar contributes adjudication-weighted forfeiture but also raises the release-backed walk option and incurs carry; the captive cells remove the release leak but retain their separately stated carry terms. Captivity does not rescue the instrument: when misconduct is punishable in place with the stake seized before any release clock runs, the leak closes, but the surviving slope is $-pC_F$, because the walk value of the balance cancels through the surplus. The stake never mints a hostage.

Corollary 1 (Turnover-funded admission conditions). *Away from the limit, positive marginal effects are cell-specific. In the leaky-exit cell they require $q > C_F + d_\theta$; equivalently, when $d_h > R$,*

$$\varphi > \varphi^* = \frac{d_h(1 - \lambda_e A) + C_F}{d_h - R}, \quad \text{feasible only if } q > C_F + R, \quad (9)$$

$$\text{ideal corner } (R = 0) : \quad \varphi > 1 - \lambda_e A + \frac{C_F}{d_h}. \quad (10)$$

The captive-exit cell requires $q > C_F$; the leaky in-place cell requires $q > C_F + (1 - s) \delta_E$; and the captive in-place cell requires $q > C_F$, in each stay cell conditional on $p > 0$.

Principal is not a self-funding source of deterrence: any positive marginal effect is financed by turnover and release surplus clearing carry plus the relevant unsanctioned recovery, and within long-lived delegation the instruments never invert dominance.

5 The filing floor and the release-clock wedge

The one instrument that can revive escrow against exit misconduct is discrimination of release speed by claim status: fast clean exits, frozen contested exits. The paper prices it.

Proposition 2 (Filing floor). *Because an unnoticed theft exits disguised as an honest departure, the thief's recovery is bounded below by $(1 - \varphi) d_h$, and the exit-theft slope is bounded above by*

$$\frac{\partial \text{cap}_x}{\partial F} \leq d_h [\lambda_e A - (1 - \varphi)] - C_F. \quad (11)$$

A necessary condition for the wedge to help at all is $\varphi > 1 - \lambda_e A$; the exact sufficient condition at the ideal corner is (10). As matches lengthen the filing requirement rises; with positive carry it exceeds one in the strict long-match limit, so no feasible filing rate produces a positive marginal escrow effect there.

False freezes are priced inside the same composites: they tax honest release value, relax the participation floor, and can help the exit constraint only in the badly filed regime. Within the model, grieving therefore worsens the wedge's own admission condition; challenger bonds, standing rules, and other anti-grieving measures are design candidates whose effects require separate modeling.

6 The regime map and the capacity prediction

Two parameters organize every case: whether unsanctioned misconduct retains the relationship (in place) or ends it (exit), and whether the unsanctioned stake is recovered or captive. Exit misconduct is deterred by the full premium and leaks stake at the thief's recovery rate; in-place misconduct is deterred by the adjudication-weighted slice of the

premium and leaks stake at the unslashed residual. The misconduct technology determines whether the relationship ends or continues, while the mechanism’s custody, filing, freeze, and slash rules determine whether stake remains leaky or captive and how recovery is priced.

The prediction concerns credible capacity, not observed misconduct rates. In the strict long-match limit, an increase in refundable principal weakly contracts the capacity frontier in all four cells, strictly in exit cells and strictly in stay cells if and only if $p > 0$. Accordingly, along the model’s strict long-match equilibrium envelope, stake size should not have a positive marginal association with credible capacity, and in exit-shaped settings the total slope is strictly negative: a larger refundable principal is not predicted to be inert; it is predicted to shrink credible capacity.

Conjecture 1 (Behavioral decomposition). Deriving a prediction for observed misconduct rates requires an additional behavioral and equilibrium model not supplied here. The natural decomposition is that the operative deterrent in long-horizon systems is the present value of the relationship: the full premium where misconduct ends the match and its p slice where it does not. In exit-misconduct settings, once the full premium is controlled for, the model-side residual capacity term associated with stake is $\varphi(d_h - R)F$, not a second scalable deterrent; it is small when φ is small or $d_h \approx R$. In in-place settings, the corresponding term is $p[d_h - (1 - s)\delta_E]F$, which is small when p is small or $d_h \approx (1 - s)\delta_E$.

A scope statement for proof-of-stake systems follows from the model’s domain and does not extend beyond it. The results apply, on the model’s primitives, to delegated staking relationships: delegator and operator, staking-as-a-service, and restaking-style delegation with exit rights and refundable principal. They do not apply to self-staked consensus participation. Within that delegated domain, the refundable fungible principal of conventional proof-of-stake delegation is not the scalable deterrent; the continuation value of remaining matched is. The author’s Secure Proof of Stake and Hybrid Secure Proof of Stake designs (Calcaterra and Kaal 2018b; Kaal 2021b, 2025) made the corresponding instrument choice, non-fungible reputation denomination and reputation-weighted rewards, without deriving the present envelope. The results do not extend to consensus-layer security, including one-shot safety attacks, where stake secures the entire chain simultaneously, attack payoffs are priced by attack-cost economics, and the token’s value is endogenous to the attack. Nothing in this paper calls into question existing proof-of-stake consensus protocols on those objects; that extension requires asset-pricing and coordination machinery outside this paper and is designated future work.

7 The market layer (conditional result)

Proposition 3 (Active-branch market equilibrium, conditional). *Under the matching-market construction maintained in this section, with rematching rate μ , burned entry fee κ , effective discount ρ , and net output function m , the active-branch reduced form is*

$$(\rho + \mu) P = m\left(\frac{P}{\theta}\right) + \mu \kappa. \quad (12)$$

Let

$$D = \rho + \mu - \frac{1}{\theta} m'\left(\frac{P}{\theta}\right). \quad (13)$$

At a differentiable solution satisfying $\mu > 0$, $D > 0$, and $P > \kappa$,

$$\frac{\partial P}{\partial \kappa} = \frac{\mu}{D} > 0, \quad \frac{\partial P}{\partial \mu} = \frac{\kappa - P}{D} < 0. \quad (14)$$

Hence fees locally raise the premium and faster rematching locally erodes it under these stated conditions. No unconditional uniqueness claim follows. Under the maintained entry construction $r\Omega = \mu \max\{P - \kappa, 0\}$ and the match-layer identity $\rho P = m(P/\theta) - r\Omega$, substitution recovers (12). Entry requires $P \geq \kappa$, hence $r\Omega \geq 0$ and $\rho P \leq m(P/\theta)$. Maintaining that $\rho P = m(P/\theta)$ has a unique positive root P_{iso} and that $\rho P - m(P/\theta) > 0$ for $P > P_{\text{iso}}$, one has $P \leq P_{\text{iso}}$ and therefore $\kappa \leq P_{\text{iso}}$; fees above P_{iso} are incompatible with entry.

What fees cannot do in this layer is enter the incentive constraint as a sanction: the fee is burned at entry, does not pad a later walk, and affects honesty only insofar as (12) raises the premium. That is a participation channel, and it follows from the fee's placement in this model, not from Theorem 1 alone. Whether an avoidable fee can additionally select among formation equilibria remains open and depends on the still-unspecified extensive-form formation protocol: under assigned continuation values with unilateral formation there is no coordination failure to select away, whereas under belief-based continuation values collapse remains a plain Nash equilibrium for every avoidable fee, and any stronger forward-induction or commitment-substitution claim remains open.

8 Design implications

Five implications follow, each conditional on the maintained model. First, size the premium, not the bond: capacity is bounded by (6) and (7), and only the premium scales

with the relationship. Second, condition the release clock on claim status, and treat filing coverage as the binding constraint it is: below the floor of Proposition 2, the wedge cannot make the marginal effect of escrow positive in the leaky-exit cell. Third, distinguish relationship retention from stake captivity: an in-place sanction changes which premium enters the constraint, while a pending-claim freeze can prevent release and move a case from leaky to captive treatment without itself converting exit misconduct into in-place misconduct. Fourth, price false freezes: within the model, grieving degrades the wedge’s own admission condition, so anti-grieving filing rules are design candidates that belong in the mechanism’s constitution rather than its appendix. Fifth, treat entry fees as participation instruments priced on the admissible interval of Proposition 3, never as a direct deterrent; any equilibrium-selection role is protocol-dependent and remains unproved.

9 Conclusion

The paper asked whether increasing refundable stake expands credible capacity when identities are disposable and relationships are long-lived. In the strict long-match limit it does not: leaky release raises walk-away value, captive regimes close that leak without producing a positive marginal principal slope, and what remains is carry. Relationship premia provide the scalable continuation-value component of deterrence, while the mechanism’s release, filing, custody, and turnover terms determine whether collateral can help at the margin away from the limit. Design by stake sizing remains design, but it is design of the wrong variable. The analysis gives exact expressions for the capacity frontier, the filing floor, the release-clock admission conditions, and the active-branch market equilibrium; equilibrium fee selection remains open pending an extensive-form formation game.

References

- Akerlof, George A., and Lawrence F. Katz. 1989. Workers’ Trust Funds and the Logic of Wage Profiles. *Quarterly Journal of Economics* 104 (3): 525-536. <https://doi.org/10.2307/2937809>
- Baker, George, Robert Gibbons, and Kevin J. Murphy. 1994. Subjective Performance Measures in Optimal Incentive Contracts. *Quarterly Journal of Economics* 109 (4): 1125-1156. <https://doi.org/10.2307/2118358>

- Becker, Gary S. 1968. Crime and Punishment: An Economic Approach. *Journal of Political Economy* 76 (2): 169-217. <https://doi.org/10.1086/259394>
- Biais, Bruno, Thomas Mariotti, Guillaume Plantin, and Jean-Charles Rochet. 2007. Dynamic Security Design: Convergence to Continuous Time and Asset Pricing Implications. *Review of Economic Studies* 74 (2): 345-390. <https://doi.org/10.1111/j.1467-937X.2007.00425.x>
- Calcaterra, Craig, Wulf A. Kaal, and Vlad Andrei. 2018. Blockchain Infrastructure for Measuring Domain Specific Reputation in Autonomous Decentralized and Anonymous Systems. https://papers.ssrn.com/sol3/papers.cfm?abstract_id=3125822
- Calcaterra, Craig, and Wulf A. Kaal. 2018b. Secure Proof of Stake Protocol. https://papers.ssrn.com/sol3/papers.cfm?abstract_id=3125827
- Carmichael, H. Lorne. 1985. Can Unemployment Be Involuntary? Comment. *American Economic Review* 75 (5): 1213-1214. <https://ideas.repec.org/a/aea/aecrev/v75y1985i5p1213-14.html>
- DeMarzo, Peter M., and Michael J. Fishman. 2007. Optimal Long-Term Financial Contracting. *Review of Financial Studies* 20 (6): 2079-2128. <https://doi.org/10.1093/rfs/hhm031>
- Friedman, Eric J., and Paul Resnick. 2001. The Social Cost of Cheap Pseudonyms. *Journal of Economics and Management Strategy* 10 (2): 173-199. <https://doi.org/10.1111/j.1430-9134.2001.00173.x>
- Friedman, James W. 1971. A Non-cooperative Equilibrium for Supergames. *Review of Economic Studies* 38 (1): 1-12. <https://doi.org/10.2307/2296617>
- Fudenberg, Drew, and David K. Levine. 1989. Reputation and Equilibrium Selection in Games with a Patient Player. *Econometrica* 57 (4): 759-778. <https://doi.org/10.2307/1913771>
- Fudenberg, Drew, and Eric Maskin. 1986. The Folk Theorem in Repeated Games with Discounting or with Incomplete Information. *Econometrica* 54 (3): 533-554. <https://doi.org/10.2307/1911307>
- Ghosh, Parikshit, and Debraj Ray. 1996. Cooperation in Community Interaction without Information Flows. *Review of Economic Studies* 63 (3): 491-519. <https://doi.org/10.2307/2297892>

- Hart, Oliver, and John Moore. 1994. A Theory of Debt Based on the Inalienability of Human Capital. *Quarterly Journal of Economics* 109 (4): 841-879. <https://doi.org/10.2307/2118350>
- Holmstrom, Bengt, and Jean Tirole. 1997. Financial Intermediation, Loanable Funds, and the Real Sector. *Quarterly Journal of Economics* 112 (3): 663-691. <https://doi.org/10.1162/003355397555316>
- Kaal, Wulf A. 2019. Decentralized Commerce: A Primer on Why Decentralized Reputation Verification Systems Are Needed. https://papers.ssrn.com/sol3/papers.cfm?abstract_id=3405401
- Kaal, Wulf A. 2021. Reputation as Capital: How DAOs Upgrade Finance. https://papers.ssrn.com/sol3/papers.cfm?abstract_id=3949098
- Kaal, Wulf A. 2021b. Hybrid Secure Proof of Stake. https://papers.ssrn.com/sol3/papers.cfm?abstract_id=3931933
- Kaal, Wulf A. 2024. AI Governance Via Web3 Reputation System. https://papers.ssrn.com/sol3/papers.cfm?abstract_id=4941807
- Kaal, Wulf A. 2025. Cryptographic Foundations and Interdisciplinary Dimensions of the Secure Proof of Stake (SPoS) Consensus. https://papers.ssrn.com/sol3/papers.cfm?abstract_id=5225296
- Kaal, Wulf A. 2026. Evolution of Domain-Specific Reputation Systems From Binary Validation to Citation-Weighted Knowledge Attribution. https://papers.ssrn.com/sol3/papers.cfm?abstract_id=6192998
- Klein, Benjamin, and Keith B. Leffler. 1981. The Role of Market Forces in Assuring Contractual Performance. *Journal of Political Economy* 89 (4): 615-641. <https://doi.org/10.1086/260996>
- Kranton, Rachel E. 1996a. The Formation of Cooperative Relationships. *Journal of Law, Economics, and Organization* 12 (1): 214-233. <https://doi.org/10.1093/oxfordjournals.jleo.a023358>
- Kranton, Rachel E. 1996b. Reciprocal Exchange: A Self-Sustaining System. *American Economic Review* 86 (4): 830-851. <https://www.jstor.org/stable/2118307>

- Kreps, David M., and Robert Wilson. 1982. Reputation and Imperfect Information. *Journal of Economic Theory* 27 (2): 253-279. [https://doi.org/10.1016/0022-0531\(82\)90030-8](https://doi.org/10.1016/0022-0531(82)90030-8)
- Levin, Jonathan. 2003. Relational Incentive Contracts. *American Economic Review* 93 (3): 835-857. <https://doi.org/10.1257/000282803322157115>
- MacLeod, W. Bentley, and James M. Malcomson. 1989. Implicit Contracts, Incentive Compatibility, and Involuntary Unemployment. *Econometrica* 57 (2): 447-480. <https://doi.org/10.2307/1912562>
- Rampini, Adriano A., and S. Viswanathan. 2010. Collateral, Risk Management, and the Distribution of Debt Capacity. *Journal of Finance* 65 (6): 2293-2322. <https://doi.org/10.1111/j.1540-6261.2010.01616.x>
- Rampini, Adriano A., and S. Viswanathan. 2013. Collateral and Capital Structure. *Journal of Financial Economics* 109 (2): 466-492. <https://doi.org/10.1016/j.jfineco.2013.03.002>
- Ray, Debraj. 2002. The Time Structure of Self-Enforcing Agreements. *Econometrica* 70 (2): 547-582. <https://doi.org/10.1111/1468-0262.00295>
- Shapiro, Carl. 1983. Premiums for High Quality Products as Returns to Reputations. *Quarterly Journal of Economics* 98 (4): 659-679. <https://doi.org/10.2307/1881782>
- Shapiro, Carl, and Joseph E. Stiglitz. 1985. Can Unemployment Be Involuntary? Reply. *American Economic Review* 75 (5): 1215-1217. <https://ideas.repec.org/a/aea/aecrev/v75y1985i5p1215-17.html>
- Telser, Lester G. 1980. A Theory of Self-Enforcing Agreements. *Journal of Business* 53 (1): 27-44. <https://doi.org/10.1086/296069>
- Thomas, Jonathan, and Tim Worrall. 1994. Foreign Direct Investment and the Risk of Expropriation. *Review of Economic Studies* 61 (1): 81-108. <https://doi.org/10.2307/2297878>