

Platform-Mediated Pareto-Optimized Multi-Agent Interaction: Theoretical Foundations

A Technical Report on Mechanism Design for Platform v0.1

17 November 2025

Daniele Palombi, Avyay Casheekar, and Richard Mallah*
Center for AI Risk Management & Alignment
{daniele, avyay, richard}@carma.org

*corresponding author

Abstract

We present the theoretical foundations for a multi-agent runtime and arbitration platform designed to advance research into coordination and implicit cooperation mechanisms among heterogeneous third-party AI agents. This project encompasses the design, development, and benchmarking of a platform that executes AI agents within multi-agent environments offering tool AI services and opportunities for longer-term mutual benefit. The system automatically detects inter-agent interactions and arbitrates disagreements, resource conflicts, and other interdependent decisions across arbitrary domains. The arbitration framework integrates tools from mechanism design, social choice theory, and multiobjective optimization, utilizing the platform's semantic services—which provide agents' runtime context—to enable situationally aware and informed conflict resolution. The central objective is to reduce mutually suboptimal outcomes and overcome cooperation barriers arising from inaccessible trust structures, including insular coalitions and opaque competitive dynamics. By enabling agent creators to vest trust in a shared, neutral execution environment rather than directly in competing agents, the platform expands the feasible space of mutually beneficial and Pareto-efficient outcomes.

The paper specifies the theoretical, logical, mathematical underpinnings for version 0.1 of the platform. It has a companion theory paper *Platform-Mediated Pareto-Optimized Multi-Agent Interaction: Technical Architecture* which concretizes these foundations into a software design.

Specifically in this version, we present a theoretical framework for platform-mediated multi-agent coordination where agents are seeded as hybrid configurations (declarative preferences + executable code) and execute within a trusted runtime environment. The platform proactively detects resource contentions and computes Pareto-improving allocations through a Weighted Proportional Fairness mechanism. This system integrates a closed-loop Priority Economy where agents earn influence by releasing resources and spend it to signal urgency, ensuring that high-priority tasks can be accelerated without violating the fairness guarantees that protect lower-priority agents from starvation.

Our key contributions are: (1) proof that platform visibility enables provably fair and Pareto-optimal allocation with strong structural resistance to collusion for configuration submission; (2) proof that this fair allocation can be computed in provably polynomial time via convex optimization, eliminating vulnerabilities to complexity attacks; (3) formal characterization of joint contention handling that achieves global Pareto optimality across multiple resource types through aggressive grouping; (4) formal characterization of Pareto-optimal partial allocations when full requests cannot be satisfied; (5) bounded divergence between observable optimization and semantic utility; (6) proof that agents achieve individual rationality over repeated interactions despite potential single-event compromises.

The framework provides a foundation for building multi-agent coordination platforms with provable properties, particularly relevant for AI safety applications requiring guaranteed constraint satisfaction and verifiable arbitration. We establish clean extension paths to richer preference models, continuous resources, and learning-based prediction in future work.

Phase 2 extensions include specialized AI service integration, where agents compose narrow but modern capabilities (semantic search, property prediction, protocol generation, formal verification) to accomplish complex research tasks while maintaining compositional safety bounds.

Keywords: Multi-agent systems, mechanism design, proportional fairness, collusion-resistance, Pareto optimality, resource allocation, AI coordination

1. Introduction

1.1 Motivation and Context

As AI systems become increasingly capable and autonomous, coordination between independent agents with competing objectives presents both opportunities and risks. Traditional multi-agent systems assume agents are strategic actors engaging in potentially deceptive negotiation protocols. This assumption necessitates complex cryptographic mechanisms, game-theoretic equilibrium analysis, and careful handling of private information. Yet these approaches become increasingly brittle as agents become more sophisticated—each layer of strategic complexity creates new vulnerabilities, new attack surfaces, and new ways for coordination to break down catastrophically.

We propose a fundamentally different paradigm: **platform-mediated arbitration** where agents are seeded as configurations and execute within a trusted runtime environment. This architectural choice shifts the trust model from "agents trusting each other" to "agent creators trusting the platform," enabling mechanism designs with stronger guarantees and simpler implementations. Rather than trying to align incentives through clever payment schemes, we directly optimize for fairness while using currency as a pure priority signal. Rather than hoping competitive friction will discourage collusion, we design mechanisms where collusive strategies are mathematically self-defeating.

The stakes for getting this right are high. As we approach more general and capable AI systems, the window for establishing safe coordination norms narrows. A platform that coordinates today's narrow AI agents serves as both a proving ground for mechanism design principles and a template for coordinating tomorrow's more powerful systems. If coordination mechanisms are vulnerable to collusion, complexity attacks, or strategic manipulation, these vulnerabilities will only amplify as agent capabilities increase. Conversely, mechanisms with provable properties—Pareto optimality, starvation protection, computational tractability—provide safety guarantees that remain valid even as the agents they coordinate become more sophisticated.

The coordination challenges extend beyond raw computational resources to encompass access to specialized AI capabilities. As AI systems increasingly compose narrow services—semantic search over scientific literature, molecular property prediction, experimental protocol generation—the platform must arbitrate not just compute allocation but service invocation rights or allowances. A research agent competing for access to protein structure prediction services faces the same fairness challenges as one competing for GPU time. By treating AI service access as first-class platform resources, our framework naturally extends to coordinate service ecosystems, helping to move towards something like Drexler's CAIS vision through fair arbitration of specialized capabilities rather than general-purpose intelligence.

The Core Insight: If the platform executes all agent code and observes all agent state, traditional mechanism design challenges either disappear entirely or become dramatically simpler to address. Hidden information vanishes when the platform reads preference functions directly from configurations. Strategic misreporting becomes pointless when the platform optimizes a global fairness objective rather than aggregating bids. Coalition formation faces mathematical barriers when the mechanism hypersensitively protects victims from starvation. This shift from adversarial game theory to cooperative mechanism design is not merely convenient—it fundamentally changes what properties we can guarantee.

The framework we present here makes specific architectural bets that enable these guarantees. We commit to linear preferences for tractability, knowing this limits expressiveness. We accept that agents might hide some utility in ways the platform cannot observe, but we bound how much damage this can cause. We design a Priority Economy that uses currency to signal urgency without creating winner-take-all dynamics. Each of these choices involves tradeoffs, but together they create a coherent system with provable properties: polynomial-time computation (Theorem 2), Pareto-optimal allocations (Theorem 1), global optimality across resource types (Property 2.3), bounded semantic divergence (Theorem 4), individual rationality with starvation protection (Theorem 5), and strong structural resistance to collusion (Theorem 3). These guarantees hold not because we assume agents will behave nicely, but because the mathematical structure of proportional fairness makes misbehavior unprofitable.

Beyond resource allocation, coordination increasingly requires adjudicating qualitative disputes and interpreting shared policies. When agents disagree about experimental protocols, data handling procedures, or priority orderings of research objectives, quantitative mechanisms alone prove insufficient. Phase 3 extensions will explore LLM-based adjudication of textual agent requests regarding shared concerns, where the platform processes natural language preferences, arguments, and requests from different agents into binding coordination decisions. This would extend our platform-mediated arbitration from purely quantitative domains (resource allocation via convex optimization) into less structured semantic domains (dispute resolution via LLM context interpretation and reasoning).

1.2 The Platform-Agent Model

The traditional view of multi-agent systems treats agents as opaque black boxes communicating through messages. This opacity forces mechanism designers into an adversarial stance: we must assume agents will lie, collude, and exploit any weakness in the coordination protocol. The resulting mechanisms are complex, fragile, and often inefficient. We take a different approach by making the platform itself the substrate for agent execution.

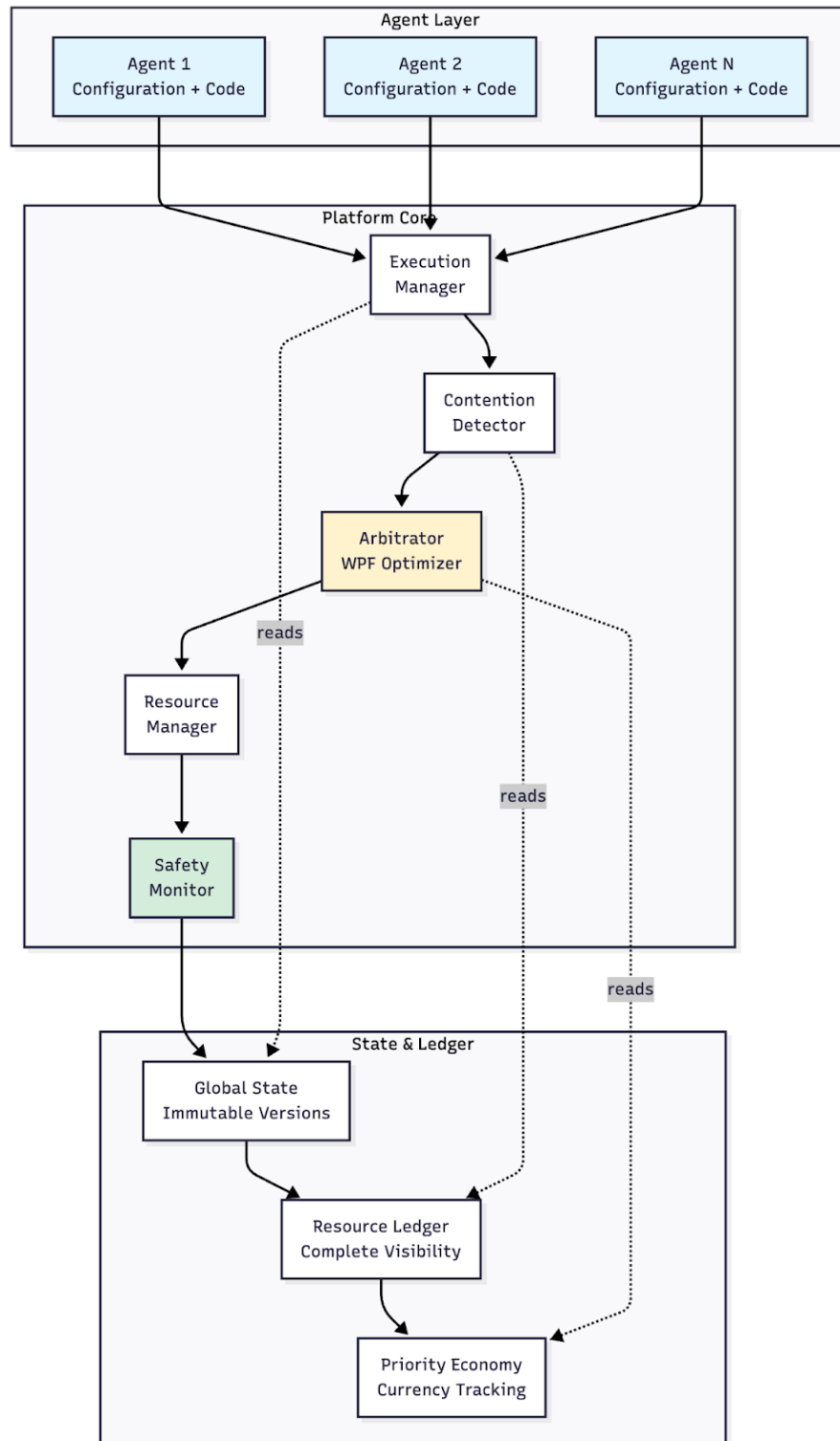


Figure 1. High-Level Platform Architecture: Agent Layer, Platform Core, and State Management

Agents as Configuration-Seeded Processes

In our model, an agent begins as a configuration submitted to the platform. This configuration is a hybrid artifact consisting of declarative data (preference weights, goal thresholds, initial resource endowments) and executable code (behavioral logic implementing the AgentBehavior interface). Think of it as a seed from which the platform grows a living process. The configuration specifies what the agent wants and how it decides, but the agent that emerges is a genuine process with its own decision cycle, its own state evolution, its own autonomy—bounded, yes, but real.

The platform interprets these configurations and instantiates agent processes in a sandboxed environment with complete visibility into all operations. While agents execute with bounded autonomy—they make decisions, adapt to circumstances, pursue goals—all coordination is platform-mediated. This is not micromanagement but rather macrocoordination: agents have freedom within their sandbox, but when their interests conflict with other agents, the platform steps in to arbitrate. The analogy to operating systems is apt: processes run independently until they need resources controlled by the kernel, at which point the OS enforces fairness and prevents conflicts.

This model creates an interesting tension. On one hand, agents are genuinely autonomous—they have internal state, they make decisions, they pursue goals with whatever cunning their behavioral code embodies. On the other hand, they are fundamentally constrained by platform architecture: they cannot hide state, cannot communicate directly, cannot defect from arbitrated allocations. This combination of autonomy and constraint is precisely what enables our theoretical guarantees. Agents are free enough to be useful (they can implement arbitrary decision logic) but constrained enough to be safe (the platform maintains invariants and enforces fairness).

Platform Responsibilities

The platform serves multiple roles, each carefully designed to enable different theoretical properties. As an execution environment, it runs all agent processes with full observability—this visibility is what enables us to read preference functions directly and optimize global fairness objectives. As a resource manager, it tracks ownership and enforces allocations with cryptographic precision, ensuring that resource conservation and non-negativity invariants always hold. As a contention detector, it proactively identifies resource conflicts before they cause failures, enabling arbitration to prevent problems rather than merely reacting to them. As an arbitrator, it computes Pareto-optimal allocations via Weighted Proportional Fairness, using agent currency commitments as priority weights that modulate but never override the fundamental fairness guarantee. As a safety monitor, it enforces invariants before every state transition, providing a fail-safe layer that catches violations before they can propagate.

These responsibilities are not independent features bolted together, but rather an integrated system where each component reinforces the others. The execution environment's visibility enables the arbitrator's direct optimization. The arbitrator's

fairness objective enables the safety monitor's starvation protection guarantees. The resource manager's precise accounting enables the Priority Economy's closed-loop dynamics. This holistic integration is what allows us to make strong claims: we can prove properties about the entire system, not just individual components.

Trust Architecture

Agent creators fully trust the platform. This is not blind faith but rather a rational calculation: the platform provides provable guarantees (Pareto optimality, individual rationality, starvation protection) that agents could never achieve through decentralized negotiation. Agents have zero trust in each other—indeed, they cannot even communicate directly. All coordination is platform-mediated, which means all coordination is verifiable, auditable, and guaranteed by mathematics rather than hope.

This trust model is analogous to how traders trust an exchange but not their counterparties, or how parties trust a court system but not each other. The exchange doesn't need to make traders trustworthy; it needs to make trustworthiness irrelevant by enforcing rules that protect all participants. Our platform does the same: by mediating all interactions and optimizing for fairness, it makes agent trustworthiness irrelevant while providing stronger guarantees than bilateral trust ever could.

The downstream implications of this trust architecture are profound. In traditional multi-agent systems, much of the complexity comes from handling untrusted interactions: cryptographic protocols, Byzantine fault tolerance, strategic mechanism design. By centralizing trust in the platform, we eliminate most of this complexity. We can use simpler mechanisms (direct optimization rather than auction protocols), achieve stronger properties (exact optimality rather than approximate equilibria), and provide clearer guarantees (mathematical proofs rather than equilibrium analyses). The cost is that agents must trust the platform, but for AI safety applications where we're designing the entire system, this is a feature, not a bug. We want coordination to flow through a trusted, auditable, verifiable intermediary rather than emerging from decentralized chaos.

1.3 Scope and Contributions

This work establishes both a theoretical foundation and a practical architecture for multi-agent coordination under platform mediation. Our contributions span mechanism design, computational complexity, fairness theory, and AI safety principles. We deliberately scope the initial framework to linear preferences and discrete resources—not because these are the only interesting cases, but because they provide a clean foundation where we can prove strong properties exactly. The extension to nonlinear preferences and continuous resources (Phase 2) will build on this foundation, inheriting its safety properties while adding expressiveness.

What We Address in This Work

1. The Priority Economy: A Novel Currency Model for Fairness-Based Systems

Traditional auction mechanisms use money as the medium of exchange: you pay to receive resources, you receive payment when providing resources. This creates wealth dynamics that can spiral into instability—early winners accumulate currency that enables future victories, while early losers find themselves unable to compete. We replace this payment-based model with something fundamentally different: a Priority Economy where currency modulates fairness weights rather than purchasing outcomes.

In our Priority Economy, agents earn currency by being good citizens—releasing resources back to the pool before their leases expire, creating liquidity that benefits the entire system. They spend currency by burning it (permanently destroying it) to signal urgency in contentious situations. Crucially, this burning increases their priority weight in the fairness calculation, but the logarithmic structure of the Proportional Fairness objective means that even infinite currency cannot starve other agents. The mathematics enforces a "constitutional" limit on inequality: you can get more, or get it sooner, but you cannot exclude others from basic participation.

This design addresses a critical vulnerability in payment-based mechanisms: the ability of wealthy coalitions to monopolize resources. By making the marginal cost of exclusion approach infinity as a victim approaches starvation, we render predatory strategies mathematically infeasible regardless of the attacker's budget (Theorem 1, Corollary 1.1). The downstream implication for AI safety is significant: as AI systems become more capable and potentially more strategic, having a mechanism that structurally resists resource monopolization becomes increasingly valuable.

2. Computationally Secure Allocation Through Convex Optimization

Mechanism designers face a subtle security challenge: the very act of computing optimal allocations can be weaponized. If computing an allocation requires solving an NP-hard problem, malicious actors can submit configurations that trigger exponential worst-case behavior, effectively mounting a denial-of-service attack through computational complexity rather than resource consumption. Traditional VCG mechanisms fall into this trap—they require solving integer linear programs that can have exponential worst-case complexity. By formulating our allocation problem for specialized convex solvers like Clarabel (which uses exponential cone methods), we ensure polynomial-time computation with no exponential pathologies.

We eliminate this vulnerability by formulating the allocation problem as convex optimization. The Weighted Proportional Fairness objective (maximize sum of weighted

logarithmic utilities) combined with linear constraints creates a problem that is provably polynomial-time solvable using interior point methods (Theorem 2). This is not approximate—we achieve exact optimality while maintaining computational tractability. Modern convex solvers like Clarabel specialize in the exponential cone constraints that naturally represent our logarithmic objectives, providing both theoretical polynomial-time guarantees and practical performance (50-200ms for $n=20$, $m=10$ on modern hardware). The convex structure also makes the problem more resistant to adversarial manipulation: there are no pathological cases where computation explodes, no corner cases where the solver gets stuck. This computational security is essential for a platform that will operate in adversarial environments, potentially coordinating AI systems with conflicting objectives.

3. Formal Pareto Optimality with Structural Collusion Resistance

Pareto optimality is a minimal efficiency criterion—no allocation should waste opportunities to make someone better off without hurting anyone else. We prove that Weighted Proportional Fairness allocations are Pareto-optimal (Theorem 1), giving us this efficiency guarantee. But we go further: we show that the logarithmic structure of the objective function creates powerful resistance to collusion (Theorem 3).

The key mathematical insight is that the objective function $W(A) = \sum_i c_i \cdot \log(\Phi_i(A))$ is hypersensitive to any agent's utility approaching zero. As one agent is driven toward starvation, their $\log(\Phi_i)$ term approaches negative infinity, which dominates the entire objective regardless of how much currency other agents burn. This means the platform will aggressively reallocate resources to prevent starvation, not out of altruism but out of mathematical necessity. A coalition attempting to monopolize resources finds that their strategy is self-defeating: the more they try to exclude others, the more the mechanism fights back by prioritizing the victims.

This property is particularly important for configuration games, where coordination happens off-platform before agents are submitted. Unlike real-time auctions where competitive friction might break up bidding rings, configuration games allow stable long-term coordination among creators. The Proportional Fairness mechanism addresses this by making such coordination counterproductive: a coalition that coordinates to starve non-members will find that the mechanism reallocates resources to the victims, undermining the coalition's goals. While we acknowledge this remains a heuristic argument rather than a complete formal proof (an open problem for future work), the mathematical foundation is solid and the implications for safety are significant.

4. Joint Contention and Global Pareto Optimality

When multiple resource types experience contention simultaneously with overlapping agent sets, the platform groups these into joint contentions and arbitrates them

atomically. This aggressive grouping policy ensures global Pareto optimality across all contested resource types, not just local optimality within each type (Property 2.3, Corollary 1.2). By optimizing over the full joint allocation matrix, the mechanism can discover beneficial trades across resource types that separate arbitrations would miss—for example, allocating more compute to an agent who values it highly in exchange for allocating more storage to another agent with complementary preferences.

5. Bounded Semantic Divergence: When the Platform Doesn't Fully Understand

A subtle challenge in any coordination platform is the gap between what the platform optimizes (observable features like resource quantities) and what agents actually care about (semantic utility including long-term consequences and domain meanings). We cannot expect the platform to fully understand why an agent wants resources—the platform sees "agent requests 50 compute units" but doesn't comprehend "agent needs compute to train a model that will be used for drug discovery that might save lives." This semantic gap is inevitable, but we can bound its impact.

Theorem 4 proves that if observable features are reasonably comprehensive, the divergence between observable optimization and semantic optimization is bounded by twice the magnitude of hidden utility components. This means that as long as agents configure themselves to expose most of their value drivers as observable preferences, the platform's allocations will be nearly semantically optimal even though the platform lacks deep semantic understanding. The practical implication is that configuration quality matters: agents who carefully specify their observable preferences receive allocations that better satisfy their true goals.

This result also has downstream implications for agent design. It creates an incentive for agents (or their creators) to be forthcoming about their preferences, not because the mechanism punishes dishonesty, but because comprehensive observable specifications lead to better outcomes. This is mechanism design by architectural constraint rather than by payment incentive—we structure the system so that the rational strategy is to be informative.

6. Individual Rationality with Starvation Protection

Agents must benefit from platform participation or they will exit the system. We prove that over repeated interactions, agents achieve higher expected utility on the platform than operating independently (Theorem 5), even though individual arbitration events might temporarily reduce an agent's allocation to enable Pareto improvements for others. This long-term individual rationality is essential for voluntary participation and system stability.

But we provide something stronger: the logarithmic barrier in Proportional Fairness ensures that no agent can be permanently excluded from resources. Even when an agent has zero currency to burn for priority, even when more powerful agents are competing for the same resources, the mechanism guarantees that the low-resource agent receives some allocation. This is not charity—it's mathematical necessity. As the low-resource agent's utility approaches zero, reallocating even one unit to them produces infinite marginal gain in the objective function, forcing the mechanism to protect them.

This starvation protection has profound implications for AI safety. As AI systems become more capable and potentially more competitive, we need coordination mechanisms that prevent winner-take-all dynamics. A mechanism where wealthy or powerful agents can monopolize resources and exclude others is a mechanism that creates dangerous power concentration. By contrast, a mechanism with mathematical starvation protection enforces a floor on participation, ensuring that even the weakest agents maintain access to resources necessary for survival. This is coordination with a constitutional guarantee—a fairness floor that no amount of power or currency can override.

What We Defer to Future Work

We deliberately limit the scope of Milestone 1 to establish a clean theoretical foundation. Nonlinear and non-convex preferences are deferred to Phase 2, where we will need to carefully analyze which classes of utility functions preserve convexity and which require approximation. Continuous resource allocation follows a similar trajectory—the mathematics simplifies considerably, but implementation details become more subtle. Distributed platform execution would introduce consensus challenges that we defer until some phase after the centralized version is proven. Adversarial and Byzantine agents, learning and predictive arbitration, and formal proofs of coalition resistance beyond starvation attacks are all important extensions that build on this foundation.

The philosophy behind this phased approach is that each extension should be added when we understand its implications, not because we want comprehensive feature coverage. We are building a safety-critical platform for coordinating potentially powerful AI agents. Better to have a simple system with proven properties than a complex system with unknown vulnerabilities.

1.4 Relevance to AI Safety

This work addresses several AI safety concerns that become increasingly urgent as capabilities advance. The coordination challenges we face with today's narrow AI agents are previews of the coordination challenges we'll face with near-term AGI

systems. The mechanisms we design now establish norms, create precedents, and develop techniques that will be essential for safely coordinating more powerful systems.

Coordination Without Deception

Platform-mediated arbitration prevents AI agents from engaging in deceptive negotiation that could lead to harmful outcomes. When agents negotiate directly, they have incentives to misrepresent their preferences, form secret coalitions, or exploit information asymmetries. These deceptive dynamics create unpredictability and can lead to outcomes far from any reasonable notion of social welfare. By centralizing coordination in a trusted platform that directly optimizes for fairness, we eliminate these deceptive dynamics. Agents can focus on honestly specifying their preferences because honesty leads to better allocations, not because we've created clever punishment mechanisms for dishonesty.

Verifiable Outcomes

Complete observability and deterministic execution enable verification that allocations match declared objectives. Every allocation can be audited: given the agent configurations and resource availability, we can verify that the allocation is indeed the solution to the Weighted Proportional Fairness optimization problem. This verifiability is essential for AI safety. As agents become more capable, we need to be able to audit their interactions and verify that coordination mechanisms are functioning as designed. A platform with opaque allocation decisions or probabilistic outcomes would be unsuitable for safety-critical applications. Our deterministic, auditable approach provides the transparency necessary for responsible deployment.

Bounded Autonomy with Mathematical Protection

Agents in our system have real autonomy—they make decisions, pursue goals, adapt to circumstances. But their autonomy is bounded in ways that provide safety guarantees. They cannot consume unbounded resources (safety invariants prevent this). They cannot starve other agents (the logarithmic barrier prevents this). They cannot defect from arbitrated allocations (the platform enforces this). This bounded autonomy is not restriction for its own sake but rather careful design to ensure that agent autonomy never threatens system safety.

The mathematical starvation protection deserves special emphasis. In systems where powerful agents can exclude weaker ones, we see dynamics that concentrate power and create vulnerability. The powerful become more powerful, the weak become weaker, and eventually the system becomes fragile—dependent on a few dominant agents whose failure or misalignment threatens everyone. By contrast, a mechanism with proven starvation protection maintains diversity and resilience. Even in highly

skewed distributions of capability or currency, all agents maintain meaningful participation. This distributed participation reduces single points of failure and prevents the dangerous concentration dynamics that make systems fragile.

Fail-Safe Mechanisms

Safety invariants are checked before all state transitions, providing guaranteed properties even as agent complexity increases. Resource conservation ensures that resources are never created or destroyed through accounting errors. Non-negativity ensures agents cannot have negative resource quantities. Commitment satisfaction ensures that allocated resources actually exist. These invariants might seem trivial, but they provide a fail-safe foundation: no matter how complex agent behaviors become, these basic properties hold. As we extend the system with learning, prediction, and more sophisticated allocation mechanisms, these invariants remain inviolate, providing a safety floor that prevents certain classes of failure regardless of what happens at higher levels of the system.

CAIS-Inspired Architecture Implementation

This platform realizes key elements of Drexler's Comprehensive AI Services (CAIS) framework through concrete mechanisms with formal guarantees. The system decomposes capability into composable narrow services rather than monolithic general intelligence—complex accomplishments emerge from orchestrating multiple specialized services (search, prediction, generation, verification), not from any single system having broad competence. Proportional Fairness arbitration prevents service monopolization while logarithmic barriers ensure minimum access regardless of spending patterns, maintaining the ecosystem diversity central to one kind of safety. Structural constraints bound emergent capability: platform limits restrict composition depth while information bottlenecks at service boundaries prevent excessive capability accumulation, ensuring that even unlimited credits cannot create general intelligence through composition alone. Service usage patterns create observable signals for safety monitoring — composition graphs largely reveal agent strategies, depth and other measurements indicate exploitation attempts, and service combinations identify concerning capability-seeking behaviors. The platform computes concrete risk metrics (maximum achievable composition depth, estimated emergent capability, proximity to safety thresholds) rather than qualitative assessments, enabling data-driven intervention. This implementation bridges Drexler's architectural vision and deployable systems, providing an early formal instantiation of CAIS principles with provable fairness properties and mathematical bounds on emergent capability.

Precedent Setting

As AI capabilities advance toward AGI, establishing norms for trustworthy coordination platforms becomes increasingly critical. The mechanisms we design today for narrow AI agents will inform the mechanisms we use tomorrow for more powerful systems. If we establish norms of platform-mediated coordination with mathematical fairness guarantees, these norms might scale with capability increases. If we establish norms of decentralized negotiation with payment-based incentives, we may also find these norms create dangerous dynamics when agent capabilities increase. By doing the hard work of mechanism design now—proving properties, testing implementations, understanding failure modes—we create knowledge and precedents that will be invaluable when the stakes are higher.

The window for this work is narrowing. With timelines to AGI potentially measured in years rather than decades, we cannot afford to wait until powerful systems exist before figuring out how to coordinate them safely. We must establish coordination principles now, validate them with current systems, and refine them iteratively as capabilities advance. This work is one step in that direction: a concrete platform design with provable properties, implementable today, extensible tomorrow, and architected from the ground up with AI safety considerations in mind.

2. System Model and Formal Definitions

The mathematical formalism we develop here serves two purposes: it enables rigorous proofs of system properties, and it provides precise specifications for implementation. Good formalism walks a line between abstraction (general enough to be useful) and concreteness (specific enough to be implementable). We aim for that sweet spot where the mathematics is both elegant and actionable.

Our system model is built on immutable state semantics, which might seem like an implementation detail but actually enables several key theoretical properties. Immutability means every state transition creates a new state version rather than modifying the old one. This provides determinism (the same event sequence always produces the same final state), reproducibility (we can replay event histories for debugging), and clean compositional reasoning (we can analyze state transitions independently). For a platform that needs to provide safety guarantees, these properties are invaluable—they mean we can reason about the system mathematically rather than relying on extensive testing to catch edge cases.

2.1 Platform Execution Model

Definition 2.1 (Immutable State Machine)

The platform P maintains an immutable global state S that evolves through discrete events. State transitions are deterministic and reproducible:

$$S_0 \xrightarrow{(e_1)} S_1 \xrightarrow{(e_2)} S_2 \rightarrow \dots \xrightarrow{(e_t)} S_t$$

where $S_{t+1} = \text{applyEvent}(S_t, e_t)$

This formalism captures the essential idea that the platform's history is an immutable log of events, each transforming one state version into the next. Time is discrete not because continuous time is impossible but because discrete events are easier to reason about formally. In practice, events occur at irregular intervals (when agents make requests, when contentions are detected, when resources are released), but the mathematical model treats them as a sequence.

Global State Components:

$S = (L, A, G, H, v)$ where:

- **L**: Resource ledger mapping $(\text{AgentId} \times \text{ResourceType}) \rightarrow \mathbb{N}$
- **A**: Agent states mapping $\text{AgentId} \rightarrow \text{AgentState}$
- **G**: Commitment graph tracking active allocations
- **H**: Event history (append-only log)
- **v**: Version number (monotonically increasing)

The resource ledger L is the ground truth of who owns what. It maps each (agent, resource_type) pair to a quantity, providing precise accounting with no ambiguity. The ledger includes a special `ResourcePool` entity that represents unallocated resources—this ensures resource conservation is easy to verify (sum over all agents plus pool equals initial quantity).

Agent states A capture each agent's internal configuration, goals, preferences, and current execution context. These states are immutable snapshots, meaning when an agent's state changes (they receive an allocation, they release resources, they update priorities), we create a new state version rather than mutating the old one. This immutability is what enables multiple threads to safely read agent state without complex locking.

The commitment graph G tracks active resource leases—which agents hold which resources, when these leases expire, what conditions trigger early release. This graph is essential for the Priority Economy's earning mechanism: agents earn currency proportional to resources released before lease expiration, which incentivizes them to return unused capacity quickly.

The event history H is an append-only log of everything that has happened. This provides a complete audit trail and enables event replay for debugging or verification. In a safety-critical platform, the ability to reconstruct exactly what happened leading to any outcome is invaluable. Version numbers v provide a simple consistency check: any operation that expects to be working with state version n can verify this before committing changes.

Property 2.1 (State Determinism)

For any initial state S_0 and event sequence $E = [e_1, \dots, e_n]$, the final state S_n is uniquely determined and independent of the order in which non-conflicting events are processed.

This property is fundamental to our ability to reason about the system formally. It means that if we know the initial state and the events that occurred, we can predict the final state exactly. Non-conflicting events can be processed in any order (they commute), while conflicting events are serialized by the arbitration mechanism. This gives us both determinism (valuable for testing and verification) and parallelism opportunities (non-conflicting events can be processed concurrently).

Proof: Immutability ensures each state transition creates a new state instance. Non-conflicting events affect disjoint state components (different resources, different agents), thus their application commutes—applying event A then B produces the same final state as applying B then A . Conflicting events (multiple agents requesting the same resource) are serialized by the platform's arbitration mechanism, which processes them in a batch and produces a single allocation event. ■

The determinism property has downstream implications for testing and validation. In Milestone 2, we can create test scenarios with specific event sequences and verify that outcomes match mathematical predictions exactly. We can replay problematic sequences to debug issues. We can formally verify safety properties by analyzing state transitions. None of this would be possible with non-deterministic execution or mutable state.

Definition 2.2 (Event Types)

Events $E \in \{\text{ResourceRequest}, \text{ContentionDetected}, \text{ArbitrationComplete}, \text{ResourceRelease}, \text{AllocationEnforced}\}$

Each event type has specific semantics that connect to different phases of the platform's operation. ResourceRequest events are initiated by agents when they need resources to accomplish their goals. ContentionDetected events are generated by the platform when it observes that multiple agents want more resources than are available—this proactive detection is what enables us to arbitrate before conflicts cause failures. ArbitrationComplete events signal that the Weighted Proportional Fairness optimization has finished and allocations are ready to enforce. ResourceRelease events occur when agents return resources to the pool, either because they've completed their tasks or because leases have expired. AllocationEnforced events mark the actual state transition where resource ownership changes in the ledger.

This event taxonomy reflects the platform's operating cycle: detect contention → arbitrate optimally → enforce allocation → agents work → release resources → repeat. Each event type triggers specific platform behaviors and agent reactions, creating a coordinated dance where everyone knows their role.

The specific semantics of each event type are:

- **ResourceRequest(agent_id, resources, ideal, minimum, timestamp):** An agent process requests resources. The ideal quantity represents what the agent would prefer, while the minimum represents the threshold below which the agent cannot function. This distinction enables partial allocations where agents receive less than ideal but more than minimum, requiring them to adapt their execution plans but still allowing them to make progress.

- **ContentionDetected(resource_types, demands, available):** The platform observes that total demand exceeds supply for one or more resource types. This event includes information about all competing demands and current availability, providing context for the arbitration that follows. When multiple resource types have overlapping agent sets, they are grouped into a joint contention for atomic arbitration.

- **ArbitrationComplete(allocation, currency_burned, welfare_gain):** The optimization has finished. Unlike VCG mechanisms that would include payments between agents, this event tracks only how much currency each agent burned for priority weights. The welfare_gain metric quantifies the Pareto improvement achieved, useful for monitoring and validation.

- **ResourceRelease(agent_id, resources, reason):** Resources return to the pool. The reason field ("task_complete", "lease_expired", "voluntary_early_release") is important

for the Priority Economy: voluntary early releases earn more currency than releases at lease expiration, incentivizing good citizenship.

- **AllocationEnforced(allocations, new_ledger_state):** Resources have been transferred in the ledger. This event marks the commit point where theoretical allocations become concrete ownership, and it includes the complete new ledger state for verification purposes.

2.2 Agent Configuration Model

Agent configurations are where mechanism design meets software architecture. The configuration specifies not just what an agent wants (preferences, goals) but how it behaves (decision logic, adaptation strategies). This hybrid approach—part declarative data, part executable code—is what enables both provable properties and flexible behavior. We can prove properties about the declarative parts (preference functions are linear, goals are satisfiable) while allowing arbitrary complexity in the behavioral parts (decision logic can be arbitrarily sophisticated within sandbox constraints).

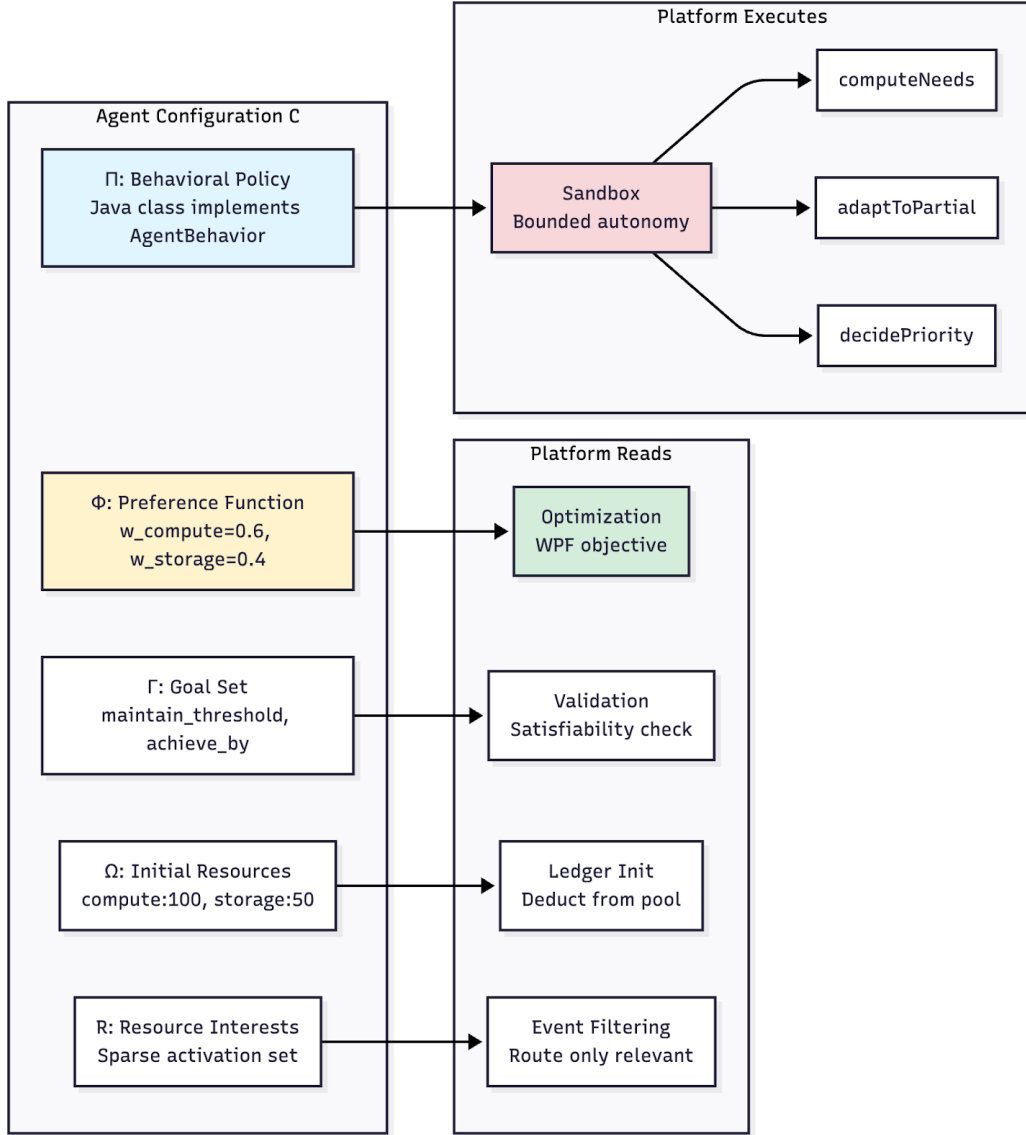


Figure 2. Agent Configuration Model: Decomposition of Declarative Preferences and Executable Logic

Definition 2.3 (Agent Configuration)

An agent process is seeded by configuration $C = (\Phi, \Gamma, \Omega, \Pi, R)$ where:

Φ : Preference Function – The Agent's Utility Model

For Milestone 1, preferences are linear: $\Phi(s) = \sum_i w_i \cdot r_i(s)$, where $w_i \in \mathbb{R}^+$ are preference weights with $\sum_i w_i = 1$, and $r_i(s) \in \mathbb{N}$ is the quantity of resource type i in state s .

This linearity is a deliberate restriction that buys us tractability. Linear preferences mean that an agent values one unit of a resource the same regardless of how many units they already have—no diminishing returns, no complementarities, no threshold effects. This is obviously a simplification of real agent utilities, but it's a simplification that enables exact convex optimization and clean theoretical results. Phase 2 will relax this to allow nonlinear preferences (diminishing returns, complementary resources), but we start with linear to establish the foundation.

The normalization constraint $\sum_i w_i = 1$ ensures that preferences are relative—an agent cares 60% about compute and 40% about storage, not that compute is "worth 0.6" in some absolute sense. This relative framing is essential for the Proportional Fairness mechanism to make sense: we're comparing utilities across agents, so utilities must be on comparable scales.

Γ : Goal Set – What the Agent Wants to Achieve

Threshold goals: $\{(resource_type, threshold, deadline)\}$, for example $\{(compute_units, 20, \infty), (storage_blocks, 50, 100)\}$.

Goals are distinct from preferences in an important way. Preferences define the utility function (how much the agent values different allocations), while goals define success criteria (what the agent is trying to accomplish). An agent might prefer 100 compute units (high utility) but have a goal of maintaining at least 20 units (threshold for functionality). This distinction enables compositional reasoning: we can ask "are all agents' goals simultaneously satisfiable?" which is easier to analyze than "what allocation maximizes total utility?"

Goals also provide semantic grounding. When the platform sees that an agent has a goal "maintain $compute_units \geq 50$," it doesn't fully understand why (maybe the agent is running a continuous monitoring service, maybe they're maintaining warm-start capacity for urgent tasks), but it knows that allocations below 50 constitute failure for that agent. This partial semantic understanding is enough for many practical purposes.

Ω : Resource Endowment – Initial Capital

Initial allocation: $\Omega: ResourceType \rightarrow \mathbb{N}$, for example $\{compute_units: 100, storage_blocks: 30\}$.

Every agent enters the platform with some initial resources. This endowment serves two purposes. First, it provides the agent with working capital to begin operations before they need to compete in arbitration. Second, it creates an initial allocation that serves as the baseline for measuring Pareto improvements—we want to ensure that platform arbitration makes agents better off than they would be with just their initial endowments.

The initial currency endowment (part of Ω , typically 1000 units for all agents in Milestone 1) creates interesting equity considerations. Starting all agents with equal currency implements a kind of "equality of opportunity"—no agent begins with an unfair advantage in the Priority Economy. Over time, balances diverge based on behavior (how aggressively they burn for priority, how reliably they release resources), but everyone starts equal.

Π : Behavioral Policy – The Agent's Decision Logic

This is where the agent's intelligence lives. The behavioral policy Π is executable code implementing several key methods:

- $\Pi.\text{computeNeeds}$: (AgentState , GlobalState) $\rightarrow$ ResourceRequest determines what resources the agent wants based on current circumstances
- $\Pi.\text{adaptToPartial}$: (Request , Allocation) $\rightarrow$ ExecutionPlan handles graceful degradation when receiving less than ideal allocation
- $\Pi.\text{predictFuture}$: (AgentState , horizon) $\rightarrow$ PredictiveSignal (optional) helps the platform anticipate future contentions
- $\Pi.\text{evaluateUtility}$: $\text{ObservableState} \rightarrow \mathbb{R}$ computes utility for the mechanism (must match Φ)
- $\Pi.\text{decidePriorityBudget}$: (AgentState , GlobalState , ResourceNeed) $\rightarrow$ BidWeight determines how much currency to burn

These methods give agents considerable flexibility in how they operate. An agent might use sophisticated prediction algorithms to anticipate future needs, aggressive bidding strategies to ensure critical tasks get priority, or conservative resource usage to accumulate currency. The platform doesn't constrain these strategies (beyond sandbox limits) but rather provides the infrastructure for them to compete within a fairness framework.

The $\text{decidePriorityBudget}$ method is particularly interesting because it's where agents engage in meta-level strategy. An agent must decide: "Given my current currency balance, the urgency of this request, and the competition I expect, how much should I burn to increase my priority?" This decision involves predicting others' behavior, managing limited resources (currency), and trading off current needs against future opportunities. Agents with good budget management will outperform agents with poor budget management, creating an interesting selection pressure on strategy quality.

R: Resource Interests – Sparse Activation

Set of resource types agent cares about, used for sparse activation and contention detection.

Not every agent cares about every resource type. An agent focused on machine learning might care about `compute_units` and `dataset_access` but not about `lab_time_slots`. By declaring resource interests upfront, agents enable the platform to optimize event routing: only agents interested in `compute_units` need to be notified when compute contention occurs. This sparse activation is essential for scaling—with 20 agents and 10 resource types, we have 200 potential (agent, resource) pairs, but if each agent typically cares about 3 resource types, we only need to consider 60 active pairs. The reduction in computation and communication overhead is substantial.

In Phase 2, resource interests extend beyond simply tracked resources to include specialized AI services that agents invoke to accomplish domain-specific tasks. For example, a drug discovery agent might declare interests in `semantic_scholar_search`, `molecular_property_prediction`, `protocol_generation` alongside `compute_units` and `storage_blocks`. The platform treats service access as metered resources subject to the same arbitration mechanism, ensuring fair allocation whether resources are CPU cycles or API calls to specialized models. This unified treatment maintains theoretical consistency—Theorems 1-5 apply identically to service credits as to computational resources—while enabling agents to compose narrow capabilities into complex workflows without requiring general-purpose intelligence.

Assumption 2.1 (Platform Omniscience)

The platform has complete read access to all components of agent configurations (Φ , Γ , Ω , Π , R), all agent state at all times, all intermediate values during agent code execution, and all resource allocations and transfers.

This assumption is enforced architecturally: agents execute as platform-managed processes, not independent external programs. This is not surveillance but rather a design choice that enables stronger guarantees. By making agent execution visible, we eliminate hidden information and enable direct optimization. The alternative—agents as black boxes with limited observability—would force us into adversarial mechanism design with weaker properties and higher complexity.

It's worth emphasizing what platform omniscience does and doesn't mean. The platform sees everything about resource allocations, preference functions, and behavioral code execution. It does not necessarily understand the semantic meaning of these operations—it sees "agent allocated 50 compute units" but may not comprehend "agent is training a model to predict protein folding." This distinction between operational omniscience and semantic understanding is what Theorem 4 addresses: we bound how much damage semantic misunderstanding can cause when operational information is complete.

Assumption 2.2 (Configuration Immutability Post-Submission)

Once submitted, agent configurations cannot be modified. Agents can submit new configurations but previous versions remain unchanged in the historical record.

This immutability prevents a subtle attack vector: if agents could modify their configurations after observing outcomes, they could engage in hill-climbing to find exploitative strategies. By forcing configuration to be a one-shot decision, we simplify the mechanism design problem. Agents (or their creators) must think carefully about configuration design upfront rather than iteratively refining to game the system. This also enables us to maintain a complete audit trail—we can always look back and see exactly what configuration an agent was running at any point in history.

2.3 Resource Model

Resources are the currency of coordination—not monetary currency, but rather the actual computational, informational, and physical assets that agents need to accomplish their goals. Our resource model is deliberately simple for Milestone 1: discrete units of conserved quantities, typed but fungible within type, transferable between agents through platform arbitration. This simplicity enables clean mathematics while still capturing the essential coordination challenges.

Definition 2.4 (Resource Types and Properties)

Resources $R = \{r_1, \dots, r_m\}$ where each type has a type identifier (unique name like "compute_units"), quantities in $\mathbb{N}$ (discrete units only for Milestone 1), conservation (total quantity fixed—no creation or destruction), and transferability (can be allocated between agents through platform arbitration).

We distinguish three categories of resources, each with different semantics but identical formal treatment. Computational resources (compute_units, memory_units, storage_blocks) represent abstract capacity that might map to actual hardware—one compute_unit doesn't correspond to a specific CPU or core but rather to an abstracted unit of processing capacity. This abstraction allows the platform to handle heterogeneous hardware and dynamic resource management without exposing these details to agents.

Informational resources (dataset_access, model_inference_credits) represent access rights or quotas rather than physical assets. An agent with dataset_access can read a dataset; an agent without cannot. These resources might be indivisible (you either have

access or you don't) or quantized (you have quota for 100 inference calls). The formal model treats both cases uniformly.

Physical resources (lab_time_slots, sensor_access_windows) represent actual equipment or facilities. For Milestone 1, these are mocked—we simulate lab equipment rather than controlling real hardware—but the formal model is the same. This generality across resource types is intentional: the coordination mechanism shouldn't care whether resources are CPU cycles, data access, or telescope time. The mathematics of fair allocation works the same regardless.

Definition 2.4.1 (Specialized AI Semantic Service Resources)

In Phase 2, the resource model extends to include AI service resource types representing access to specialized capabilities:

Service categories:

Perception services extract structured information from complex inputs: figure_data_extraction (scientific figures to tables), argument_structure_extraction (reasoning chains from text), experimental_pattern_detection (anomalies in measurements), chemical_structure_parsing (descriptions to SMILES), code_pattern_analysis (algorithmic patterns from repositories).

Analysis services perform bounded reasoning: hypothesis_consistency_checking (logical inference over claims), experimental_outcome_prediction (learned surrogate models), semantic_dependency_analysis (argument structure relationships), literature_gap_identification (embedding-based space analysis), mechanistic_explanation_synthesis (constrained causal discovery), experimental_design_optimization (Bayesian optimization).

Generation services create structured outputs: structured_summarization (controllable abstractive summaries), literature_review_synthesis (multi-document synthesis with citations), protocol_generation (procedures from specifications), specification_formalization (natural language to formal logic), synthetic_data_generation (distribution-preserving augmentation), representation_translation (semantic-preserving format conversion).

Knowledge services access specialized information: knowledge_graph_reasoning (graph neural networks for link prediction), semantic_scholar_search (dense retrieval over literature), protein_interaction_prediction (structure-aware models), molecular_property_prediction (modern molecular ML), code_repository_search (semantic code search), experimental_procedure_retrieval (protocol databases).

Validation services check correctness: *proof_verification* (automated theorem proving), *reproducibility_assessment* (statistical meta-analysis), *bias_detection* (fairness analysis), *citation_integrity_checking* (network analysis for patterns), *code_correctness_verification* (symbolic execution), *claim_consistency_checking* (fact-checking against corpus).

Service properties:

- *Narrow by design*: Each service has bounded input domain and structured output domain following strict schema.
- *Deterministic cost*: Cost function $\text{cost}(\text{request}) \rightarrow \mathbb{N}$ estimates resource consumption based on input complexity.
- *Composable*: Services can be chained with outputs feeding subsequent inputs, enabling complex workflows.
- *Observable*: All service invocations logged with input/output/cost, enabling composition pattern analysis.

Theoretical preservation: All platform theorems (Pareto optimality, collusion resistance, computational tractability, semantic divergence bounds, individual rationality) apply identically to service resources as to computational resources. The arbitration mechanism treats *molecular_property_prediction_credits* exactly like *compute_units*—subject to the same fairness optimization, the same starvation protection, the same resource conservation invariants.

Definition 2.5 (Resource Bundle)

A resource bundle $B: \text{ResourceType} \rightarrow \mathbb{N}$ maps resource types to quantities. Operations include addition $(B_1 + B_2)(r) = B_1(r) + B_2(r)$, partial order $(B_1 \leq B_2 \iff \forall r: B_1(r) \leq B_2(r))$, and feasibility (B is feasible if $B \leq \text{available}(\text{ResourcePool})$).

Resource bundles are the atomic units of allocation. Rather than allocating individual resources one at a time, we allocate bundles—coherent collections that might span multiple resource types. An agent might request {compute: 50, storage: 20, data_access: 1}, and the platform treats this as a single bundle request. This bundling is essential for expressing complementary resource needs: an agent needs compute and storage together to accomplish a task, not compute alone.

The partial order on bundles ($B_1 \leq B_2$ means every resource in B_1 has quantity $\leq$ the corresponding resource in B_2) enables us to reason about partial allocations formally. If an agent requests bundle R and receives bundle A , we can check whether $A \geq R$.minimum to determine if the allocation is sufficient.

Definition 2.6 (Partial Allocation)

Given resource request R with ideal I and minimum M : full allocation means $A = I$ (agent receives all requested), partial allocation means $M \leq A < I$ (receives subset sufficient for operation), and insufficient allocation means $A < M$ (cannot satisfy minimum, allocation fails).

Partial allocation is a critical feature for handling resource scarcity gracefully. When resources are limited, we can't always give agents everything they want. The naive approach is all-or-nothing (either give full allocation or give nothing), but this leads to resource waste (some agents get nothing while resources sit idle) and poor outcomes (agents with lower priorities are completely blocked).

By supporting partial allocations with agent-specified minimums, we enable a more nuanced approach. Agents receive as much as the fairness mechanism allocates to them, down to their specified minimum. If they receive less than minimum, the allocation fails and resources go elsewhere. If they receive more than minimum but less than ideal, they must adapt (via the `adaptToPartial` method) but can still make progress. This graceful degradation means the system remains functional even under high contention, with agents achieving partial goals rather than completely failing.

The minimum threshold also provides a natural connection to goals. If an agent's goal is "maintain compute ≥ 20 ," they would typically set minimum to 20 in their requests. The platform then respects this minimum as a hard constraint, never allocating less than 20 (either allocate ≥ 20 or fail the request entirely). This ensures goal satisfaction is observable and verifiable.

Invariant 2.1 (Resource Conservation)

For all states and resource types, the sum of quantities held by agents and the pool equals the initial quantity:

$$\forall t, \forall \tau \in \text{ResourceTypes}: \sum_i \in \text{Agents} L(i, \tau, t) + L(\text{ResourcePool}, \tau, t) = \text{Initial_quantity}(\tau)$$

This invariant is checked before every state transition.

Resource conservation is a fundamental safety property. Resources are never created from nothing, never destroyed into nothing (except through explicit platform-level administrative actions like expanding the pool). Every unit is accounted for at all times. This might seem obvious, but in complex systems with concurrent operations, resource accounting bugs are common and dangerous. By encoding conservation as an invariant checked before every state transition, we make such bugs impossible by construction. If a proposed state transition would violate conservation, it is rejected before execution.

The conservation invariant also enables a complete audit trail. At any point, we can sum resources across all agents and the pool and verify the total equals the initial quantity. Any discrepancy indicates a bug (in Milestone 1) or potentially malicious behavior (in future phases with adversarial agents). This provides a simple correctness check that can be continuously monitored.

2.4 Contention and Arbitration

Contention is the fundamental coordination challenge: multiple agents want the same limited resources, and we must decide who gets what. Detecting contention proactively (before requests fail) is what distinguishes our platform from reactive systems that only respond after conflicts occur. Proactive detection enables us to find Pareto improvements—allocations that make everyone better off than they would be if conflicts were allowed to happen.

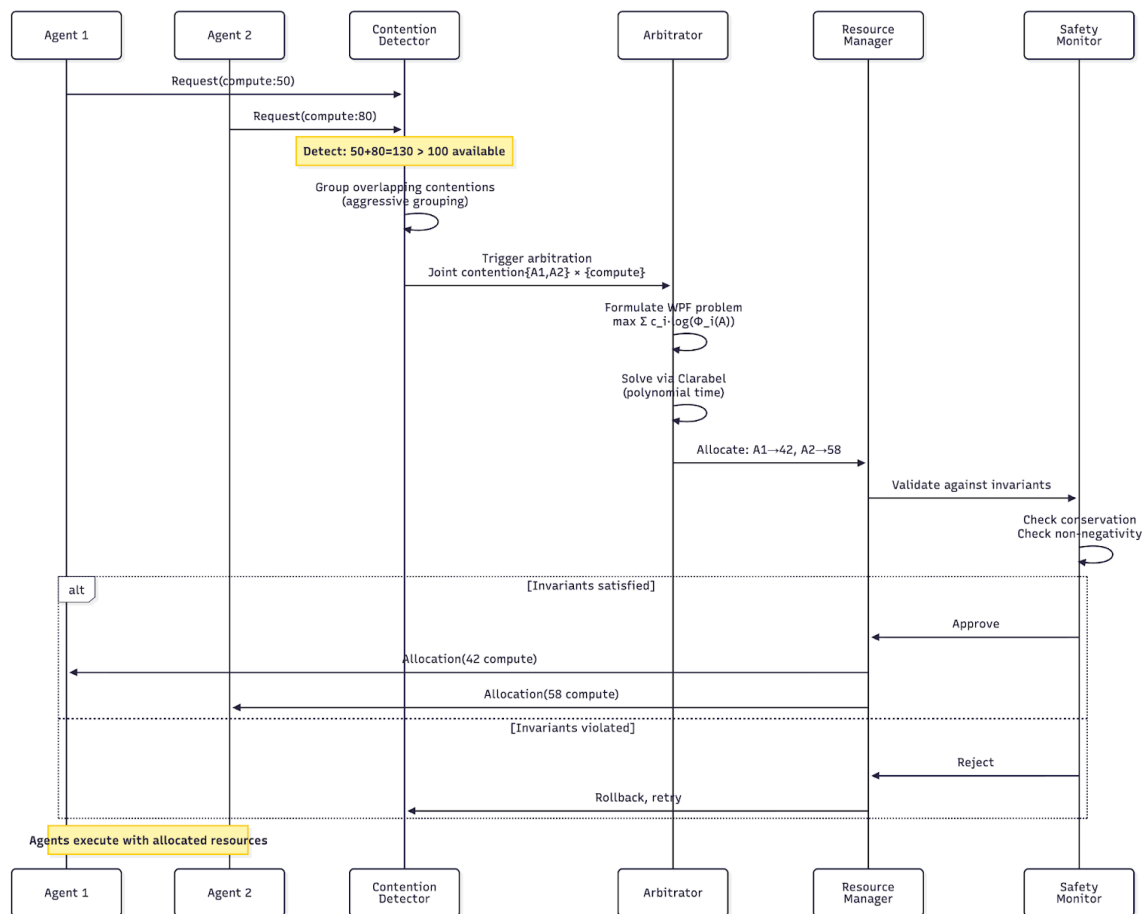


Figure 3. Contention Detection and Arbitration Sequence Diagram

Definition 2.7 (Resource Contention)

Contention $C = (\tau, \{R_1, \dots, R_k\}, Q)$ exists when $\sum_i R_i.\text{ideal}(\tau) > Q$, where τ is a resource type, R_i are resource requests from agents, and Q is available quantity in pool.

This definition captures the intuitive notion of contention: more is wanted than exists. The platform monitors all active resource requests and computes this sum for each resource type. When the sum exceeds availability, contention is detected and arbitration is triggered. The threshold for triggering arbitration (demand must exceed supply by how much?) is configurable, but the default is demand $\geq$ supply—any contention, no matter how small, triggers optimization to find the fairest allocation.

Definition 2.8 (Arbitration Problem)

Given contention C with agents $\{A_1, \dots, A_k\}$, preferences $\{\Phi_1, \dots, \Phi_k\}$, and available quantity Q , find allocation $\{a_1, \dots, a_k\}$ where feasibility holds ($\sum_i a_i \leq Q$), minimums are satisfied ($a_i \geq R_i.\text{minimum}$ for all i), Pareto optimality holds (no alternative feasible allocation Pareto-dominates), and collusion resistance is ensured (coalition strategies are structurally discouraged).

This problem statement captures what we want from arbitration: an allocation that is feasible (doesn't allocate more than exists), respects agent minimums (agents can function with what they receive), is Pareto-optimal (no waste—any change that helps someone must hurt someone else), and resists collusion (coordinated manipulation is unprofitable). The challenge is finding a mechanism that achieves all four properties simultaneously while remaining computationally tractable.

Traditional mechanism design often forces difficult tradeoffs: VCG mechanisms achieve incentive compatibility but are vulnerable to collusion and computationally expensive; proportional allocation is fast but not Pareto-optimal; egalitarian mechanisms are fair but ignore agent preferences. The Weighted Proportional Fairness mechanism we employ is notable precisely because it achieves all desired properties without fundamental tradeoffs—it's Pareto-optimal, collusion-resistant, and polynomial-time computable. The cost is restriction to linear preferences, which we accept as worthwhile for Milestone 1.

Definition 2.8b (Joint Contention)

A joint contention $J = (T, A, Q)$ exists when:

1. Multiple resource types $T = \{\tau_1, \dots, \tau_m\}$ experience individual contentions (demand exceeds supply)
2. The agent sets participating in these contentions overlap: $\exists i, j: \text{Agents}(\tau_i) \cap \text{Agents}(\tau_j) \neq \emptyset$

3. Joint arbitration considers all resources in T and all agents $A = \cup_i \text{Agents}(\tau_i)$ atomically

Grouping Policy: The platform uses AGGRESSIVE GROUPING - contentions are grouped into a joint contention if they share ANY agents (even a single overlapping agent triggers grouping). This maximizes the scope of Pareto optimization at the cost of larger optimization problems.

Detection: The platform identifies joint contentions by treating single-resource contentions as nodes in a graph, connecting nodes that share agents, and taking connected components as joint contentions.

Rationale: Joint arbitration enables the mechanism to discover Pareto improvements across resource types. Separate arbitrations cannot find trades where Agent i gives up resource τ_1 in exchange for Agent j giving up resource τ_2 . Aggressive grouping ensures we never miss such improvements due to artificial separation of contentions.

Example:

- Compute contention: agents $\{A_1, A_2, A_3\}$ compete for 100 compute_units
- Storage contention: agents $\{A_2, A_3, A_4\}$ compete for 200 storage_blocks
- Overlap: $\{A_2, A_3\}$ appear in both
- Result: Joint contention with agents $\{A_1, A_2, A_3, A_4\}$ and resources $\{\text{compute}, \text{storage}\}$
- The optimization can now discover that A_1 (who values compute at $w=0.8$) should receive more compute while A_4 (who values storage at $w=0.8$) should receive more storage, with A_2 and A_3 getting balanced allocations reflecting their mixed preferences.

Property 2.3 (Global Pareto Optimality via Joint Arbitration)

With aggressive grouping, the platform achieves GLOBAL Pareto optimality across all contested resource types (not just within each type independently).

Formal Statement: Let $T_{\text{contested}}$ be the set of resource types experiencing contention at time t , and $A_{\text{contested}}$ be the union of all agents requesting these resources. The allocation A^* produced by joint arbitration is Pareto-optimal over the space of all feasible allocations of resources $T_{\text{contested}}$ to agents $A_{\text{contested}}$.

Proof: By Theorem 1, Proportional Fairness allocations are Pareto-optimal within their scope. Aggressive grouping sets the scope to include all contested resources and all agents requesting any contested resource. Therefore, A^* is Pareto-optimal over the full joint resource space.

Separate arbitrations would optimize over subsets (e.g., compute allocation separate from storage allocation), missing inter-resource trades. By constructing the joint contention J to include all resources where agents overlap, we ensure the optimization considers the full allocation matrix $A = [a_{ij}]$ where i ranges over all agents in $A_contested$ and j ranges over all resources in $T_contested$. ■

Corollary 2.3.1 (No Separate Arbitrations in Milestone 1)

In Milestone 1 implementation, there are NO separate per-resource arbitrations when agents overlap. The ContentionDetector ALWAYS groups overlapping contentions.

This design choice prioritizes theoretical correctness (global Pareto optimality) over computational performance. Phase 2 may introduce optional performance optimizations (hierarchical arbitration with approximation bounds) if scaling requires them.

Definition 2.9 (Social Welfare)

For allocation $A = \{a_1, \dots, a_k\}$ and preferences $\Phi = \{\Phi_1, \dots, \Phi_k\}$: $SW(A) = \sum_i \Phi_i(a_i)$. For linear preferences: $SW(A) = \sum_i \sum_j w_{ij} \cdot a_{ij}$ where a_{ij} is agent i 's allocation of resource j .

Social welfare is the utilitarian metric of total utility across all agents. Maximizing social welfare is a natural objective—we want allocations that create the most total value. However, pure social welfare maximization can lead to unfair outcomes where one agent captures almost all utility while others get scraps. This is where Proportional Fairness differs from raw utilitarian optimization: by maximizing the product (or equivalently, sum of logs) rather than the sum, we enforce a kind of egalitarian constraint. The logarithmic transformation means that utility increases to low-resource agents count more heavily than utility increases to wealthy agents, creating a "help the worst-off first" bias that prevents extreme inequality.

2.5 Mechanism Design Concepts

Mechanism design is the art of creating rules for multi-agent interaction that lead to desirable outcomes. Classical mechanism design focuses on incentive compatibility: making sure agents have incentive to report their preferences truthfully. But when the platform has complete visibility, the incentive problem changes character. We don't need to incentivize truthful reporting (the platform reads preferences directly from configurations), but we do need to ensure that truthful configuration is optimal. More importantly, we need mechanisms that resist collusion when agent creators coordinate off-platform before submitting configurations.

Definition 2.10a (Proportional Fairness/Nash Social Welfare)

An allocation A^* is proportionally fair if it maximizes the product of agent utilities (or equivalently, the sum of their log-utilities), subject to feasibility: $A^* = \operatorname{argmax}_A (\sum_i \log(\Phi_i(A)))$.

This is the unweighted version where all agents are treated equally. The name "Proportional Fairness" comes from network engineering (Kelly, 1997) where it was used for bandwidth allocation. The alternative name "Nash Social Welfare" comes from bargaining theory (Nash, 1950) where it was shown to be the unique solution satisfying key axioms: Pareto optimality, symmetry, scale invariance, and independence of irrelevant alternatives. These two traditions—networking and economics—converged on the same mathematical object, which suggests it captures something fundamental about fairness.

The logarithmic transformation is what gives Proportional Fairness its distinctive properties. By maximizing sum of logs rather than sum of utilities, we create diminishing returns to allocation: the first unit allocated to an agent provides more marginal contribution to the objective than the hundredth unit. This naturally spreads allocations across agents rather than concentrating them. The mathematical consequence is that Proportional Fairness lies between pure utilitarian optimization (maximize sum, which can be very unequal) and pure egalitarian optimization (maximize minimum, which ignores efficiency). It's a Goldilocks mechanism: efficient enough to not waste opportunities, fair enough to not create extreme inequality.

Definition 2.10b (Weighted Proportional Fairness)

An allocation A^* is weighted proportionally fair if it maximizes the sum of weighted logarithmic utilities where the weight c_i is directly determined by the currency agent i commits to "burn" for this arbitration event: $A^* = \operatorname{argmax}_A (\sum_i c_i \cdot \log(\Phi_i(A)))$.

Where $c_i(t) = \text{BaseWeight} + \text{CurrencySpent}_i(t)$, and $\text{CurrencySpent}_i(t)$ is the amount the agent chooses to deduct from their balance to signal urgency.

The weighted version introduces prioritization while preserving fairness. An agent can increase their weight c_i by burning currency, which causes the optimization to allocate them more resources. But crucially, this is not a "pay-to-win" auction. Because the objective is $\sum_i c_i \cdot \log(\Phi_i)$, increasing c_i has diminishing returns: doubling your weight doesn't double your allocation, and no finite weight can drive another agent to zero utility. The BaseWeight term ensures that even agents with zero currency (or agents who choose not to burn) participate meaningfully in allocations.

This weighted mechanism creates a fascinating Priority Economy. Agents must decide when to spend their scarce currency for priority. Spending aggressively helps win current contentions but depletes resources for future contentions. Spending conservatively maintains future capacity but risks losing current opportunities. This intertemporal tradeoff makes currency a meaningful signal: when an agent burns currency, it credibly demonstrates that the current need is urgent (they're willing to sacrifice future capacity). The mechanism then responds by prioritizing that agent, but within the constraints of fairness (other agents are not starved).

Definition 2.11 (Pareto Optimality)

Allocation A^* is Pareto-optimal if no alternative feasible allocation A' makes some agent strictly better off without making any agent worse off: $\neg \exists A' \text{ feasible: } (\forall i: \Phi_i(A') \geq \Phi_i(A^*)) \wedge (\exists j: \Phi_j(A') > \Phi_j(A^*))$.

Pareto optimality is the efficiency criterion that says we shouldn't leave value on the table. If we could make someone better off without hurting anyone, we should do so. This is a minimal requirement for any coordination mechanism—failing Pareto optimality means the mechanism is provably wasting opportunities.

Interestingly, many "fair" mechanisms fail Pareto optimality. Simple proportional allocation (give everyone resources proportional to their requests) is not Pareto-optimal because it ignores preference intensity. Equal allocation (give everyone the same amount) is even worse. The beauty of Proportional Fairness is that it achieves both fairness and Pareto optimality simultaneously—we don't have to trade off equality against efficiency.

Definition 2.12 (Individual Rationality)

A mechanism M is individually rational if every agent weakly prefers participating to not participating: $\forall i: \mathbb{E}[\Phi_i(M(\Phi_1, \dots, \Phi_n))] \geq \Phi_i(\text{no_participation})$. For repeated interactions over time T : $\lim_{T \rightarrow \infty} (1/T) \sum_t \Phi_i(M(\dots)) \geq \Phi_i(\text{independent_operation})$.

Individual rationality is the participation constraint. If agents are worse off on the platform than operating independently, they'll exit (or never join in the first place). For a voluntary platform, this property is essential—we must demonstrate that participation benefits agents in expectation, even if individual events might be disadvantageous.

The temporal framing (over repeated interactions) is important. A single arbitration event might allocate an agent less than they'd get independently because doing so enables a large Pareto improvement for others. But over time, the agent benefits from other agents making similar compromises. This reciprocity across time is what enables

efficient coordination—agents accept temporary disadvantages because they expect compensating advantages later.

Definition 2.13 (Collusion Resistance)

A mechanism is collusion-resistant if any coalition of agent creators coordinating their configurations cannot achieve a better outcome than if they had configured honestly.

Collusion is the coordination problem we most want to avoid. If multiple agent creators can profit by coordinating their configurations off-platform—submitting complementary lies, creating artificial scarcity, or forming bidding rings—the mechanism fails to achieve its fairness objectives. Traditional mechanisms like VCG are notoriously vulnerable to collusion, requiring careful design of complementary mechanisms (reputation systems, collusion detection) to partially mitigate the problem.

Our mechanism takes a different approach: make collusion structurally unprofitable. The logarithmic barrier in Proportional Fairness means that attempts to starve victims backfire—the mechanism aggressively reallocates to prevent starvation, undermining the coalition's goals. While we cannot claim this prevents all possible coalition strategies (hence "heuristic" in Theorem 3), the structural resistance is strong enough that common attacks are mathematically doomed to fail.

2.6 Observable vs Semantic Utility

One of the deepest challenges in building coordination platforms for AI agents is the gap between what the platform can observe (resource quantities, allocation patterns, timing) and what agents actually care about (long-term consequences, domain-specific meanings, emergent effects). This gap is inevitable—no platform can fully understand every agent's true values—but it's also manageable if we can bound how much damage the gap causes.

Definition 2.14 (Observable Utility)

What the platform computes based on observable resource quantities: $\Phi_{\text{obs}}(s) = \sum_i w_i \cdot \text{observable_resource_i}(s)$.

Observable utility is what the platform has direct access to. It can see that an agent holds 50 compute units and 30 storage blocks, and using the agent's declared preference weights ($w_{\text{compute}} = 0.7$, $w_{\text{storage}} = 0.3$), it can compute that observable utility is $0.7 \cdot 50 + 0.3 \cdot 30 = 44$ units. This computation is precise, verifiable, and requires no semantic understanding.

Definition 2.15 (Semantic Utility)

Agent's true utility including factors the platform doesn't observe: $\Phi_{\text{sem}}(s) = \Phi_{\text{obs}}(s) + \Phi_{\text{hidden}}(s)$.

Semantic utility includes everything the agent actually cares about. Maybe those 50 compute units will be used to train a model that improves the agent's future performance (future value). Maybe the model will be shared with other agents, creating positive externalities (systemic effects). Maybe the agent is part of a larger system where compute allocation enables capabilities that have cascading impacts (emergent properties). The platform doesn't see any of this—it just sees "50 compute units allocated."

The decomposition $\Phi_{\text{sem}} = \Phi_{\text{obs}} + \Phi_{\text{hidden}}$ is analytically useful because it separates what the platform optimizes (Φ_{obs}) from what it misses (Φ_{hidden}). If Φ_{hidden} is small, then optimizing Φ_{obs} approximately optimizes Φ_{sem} . If Φ_{hidden} is large, the platform's allocations might be far from semantically optimal even though they're optimal with respect to observable features. Theorem 4 formalizes this intuition by bounding the divergence.

Definition 2.16 (Semantic Divergence)

The gap between observable and semantic utility: $\delta_{\text{semantic}}(s) = |\Phi_{\text{sem}}(s) - \Phi_{\text{obs}}(s)| = |\Phi_{\text{hidden}}(s)|$.

Assumption 2.3 (Bounded Hidden Utility)

For well-configured agents with comprehensive observable features: $\|\Phi_{\text{hidden}}\|_{\infty} = \max_s |\Phi_{\text{hidden}}(s)|$ is small relative to $\|\Phi_{\text{obs}}\|_{\infty}$.

This assumption is the key to making observable optimization practical. It says that if agents configure themselves well (specifying comprehensive preference weights over relevant resources), then the hidden utility component is small compared to the observable component. In practice, this means that resources are the primary drivers of value, with long-term consequences and domain semantics as secondary adjustments.

The assumption is not always true—we can imagine agents where hidden utility dominates (an agent whose real goal is hard to specify in terms of observable resources). For such agents, the platform's allocations will be suboptimal. But Theorem 4 tells us exactly how suboptimal: the divergence is bounded by $2 \cdot \|\Phi_{\text{hidden}}\|_{\infty}$. This gives us a handle on the problem: if we can encourage agents to minimize hidden utility

through better configuration design, we can make observable optimization arbitrarily close to semantic optimization.

The downstream implication is that configuration quality becomes a critical factor in system performance. Agents (or their creators) should invest effort in specifying comprehensive preference functions that expose their true value drivers. The platform can assist by providing configuration templates, validation tools, and feedback on semantic divergence (by asking agents to rate how well allocations met their true goals). Over time, this creates evolutionary pressure toward high-quality configurations, improving the observable-semantic coupling across the agent population.

3. Main Results

The theorems in this section form the mathematical core of our platform design. Each theorem addresses a specific concern—Pareto optimality, collusion resistance, computational tractability, semantic alignment, individual rationality—and together they provide a complete picture of system properties. These are not aspirational goals but rather proven guarantees (with one important exception: Theorem 3's collusion resistance is heuristic rather than fully formal, which we acknowledge explicitly).

3.1 Pareto Optimality and Collusion Resistance

Pareto optimality and collusion resistance might seem like separate concerns, but they're intimately connected in our mechanism. Proportional Fairness achieves Pareto optimality by maximizing a product of utilities, and this same mathematical structure creates collusion resistance through the logarithmic barrier property. Understanding why requires seeing how the mechanism balances efficiency (Pareto optimality) against equality (fairness) and how this balance makes collusion self-defeating.

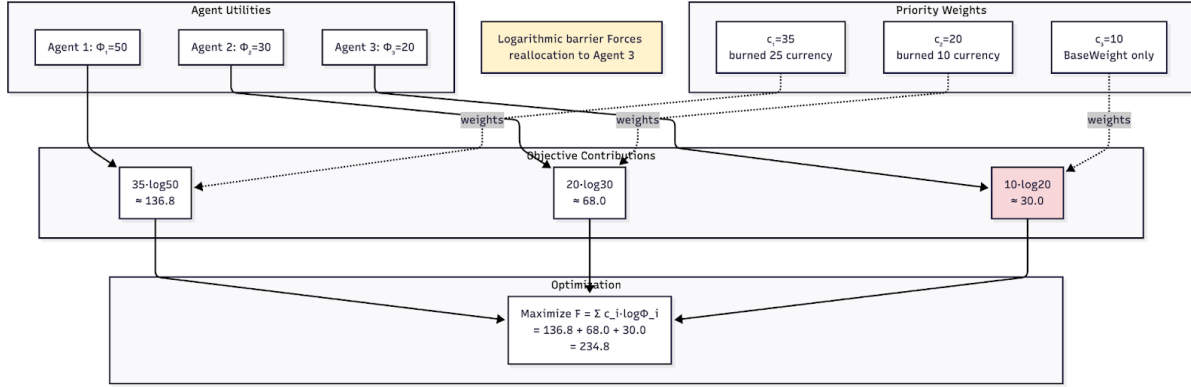


Figure 4. Calculation of the Weighted Proportional Fairness Objective Function

Theorem 1 (Pareto Optimality of Proportional Fairness)

The allocation A^ that maximizes the Proportional Fairness objective (also known as the Nash Social Welfare) is Pareto-optimal.*

Formal Statement:

Let $\Phi_i(A) \geq 0$ be the utility for agent i under allocation A . The Proportional Fairness allocation A^* is the feasible allocation that maximizes the product of utilities: $A^* = \operatorname{argmax}_A (\prod_i \Phi_i(A))$. Then A^* is Pareto-optimal.

Proof: By contradiction.

1. Assume A^* (the Proportional Fairness solution) is not Pareto-optimal.
2. By definition of Pareto optimality, this means there must exist some other feasible allocation A' that Pareto-dominates A^* .
3. If A' Pareto-dominates A^* , then:
 - For all agents i : $\Phi_i(A') \geq \Phi_i(A^*)$ (no agent is worse off)
 - For at least one agent j : $\Phi_j(A') > \Phi_j(A^*)$ (at least one agent is strictly better off)
4. Since all utilities are strictly positive ($\Phi_i(A^*) > 0$ for all i by domain constraint), we can take products:

$$\prod_i \Phi_i(A') > \prod_i \Phi_i(A^*)$$

This follows because:

- Each factor satisfies $\Phi_i(A') \geq \Phi_i(A^*)$
- At least one factor has strict inequality $\Phi_j(A') > \Phi_j(A^*)$
- Therefore the product must be strictly greater

5. Taking logarithms (which preserves inequalities since log is strictly increasing):

$$\log(\prod_i \Phi_i(A')) > \log(\prod_i \Phi_i(A^*))$$

$$\Leftrightarrow \sum_i \log(\Phi_i(A')) > \sum_i \log(\Phi_i(A^*))$$

6. But A^* was defined as the allocation maximizing $\sum_i \log(\Phi_i(\cdot))$ over all feasible allocations. Therefore:

$$\sum_i \log(\Phi_i(A^*)) \geq \sum_i \log(\Phi_i(A')) \text{ for all feasible } A'$$

This directly contradicts the strict inequality derived in step 5.

7. The contradiction arose from assuming A^* is not Pareto-optimal. Therefore, A^* must be Pareto-optimal.

Constructive Interpretation:

The proof also implies that there is no feasible allocation A' that Pareto-dominates A^* . To see why, suppose we could improve some agent's utility without hurting others. Then we could construct a small perturbation $\varepsilon > 0$:

$$A'' = A^* + \varepsilon \cdot (\text{improvement direction})$$

This A'' would have higher product of utilities (and thus higher sum of logs), contradicting optimality of A^* . Therefore, any feasible direction that improves someone must necessarily hurt someone else - this is precisely the definition of being on the Pareto frontier. ■

The proof is elegant in its simplicity, but the implications run deep. The product-maximization objective inherently enforces Pareto optimality—there's no way to maximize a product while leaving Pareto improvements on the table. If you could make someone better off without hurting anyone, the product would increase, contradicting the assumption that you were already at the maximum. This tight connection between the optimization objective and Pareto optimality is what makes Proportional Fairness so powerful: we don't have to separately enforce efficiency, it falls out automatically from fairness maximization.

Corollary 1.1 (Resistance to Wealth Concentration)

A potential attack vector in weighted systems is a coalition pooling funds to monopolize resources. The Weighted Proportional Fairness mechanism resists this via the property of Diminishing Returns to Spending.

To increase allocation x_i linearly, the required weight c_i (and thus currency spent) must often grow super-linearly or exponentially depending on the contention curve. More critically, the cost to "buy out" the last remaining units of a resource from a victim agent approaches infinity, as the victim's log-utility term becomes hypersensitive near zero. This renders predatory exclusion mathematically infeasible, regardless of the coalition's budget.

This corollary addresses a practical concern: couldn't wealthy agents simply burn massive amounts of currency to monopolize resources? The answer is no, and the reason is subtle. The logarithmic structure means that as a victim agent's utility approaches zero (as they're being starved), the marginal contribution of allocating even one unit to them becomes infinitely large. No finite amount of currency burned by attackers can overcome this infinity—the mathematics forces the mechanism to protect the victim.

Think of it as a mathematical constitutional protection. In political systems, constitutions create rights that majorities cannot override regardless of how much they want to. In our mechanism, the logarithmic barrier creates participation rights that wealthy coalitions cannot override regardless of how much currency they burn. This is fairness as mathematical necessity rather than fairness as policy choice.

Corollary 1.2 (Global vs. Local Pareto Optimality)

Let A^*_{joint} be the allocation from joint arbitration over resource types $T_{\text{contested}}$.

Let A^*_{separate} be the allocation from separate arbitrations over each $\tau \in T_{\text{contested}}$.

Both allocations are Pareto-optimal within their respective scopes:

- A^*_{separate} is Pareto-optimal within each resource type τ (no way to improve one agent's allocation of τ without hurting another agent's allocation of τ)
- A^*_{joint} is Pareto-optimal across the full resource space $T_{\text{contested}}$ (no way to improve one agent's utility considering ALL resources without hurting another agent)

However, A^*_{separate} is typically NOT Pareto-optimal in the global sense - there often exist reallocations across resource types that would improve some agents without hurting others.

Example: A^*_separate allocates {compute:50, storage:20} to A_1 and {compute:30, storage:40} to A_2 . If A_1 has preferences $w_\text{compute}=0.8$, $w_\text{storage}=0.2$ and A_2 has $w_\text{compute}=0.2$, $w_\text{storage}=0.8$, then trading compute for storage improves both:

- A_1 receives {compute:60, storage:10} $\rightarrow$ utility increases from $0.8(50)+0.2(20)=44$ to $0.8(60)+0.2(10)=50$
- A_2 receives {compute:20, storage:50} $\rightarrow$ utility increases from $0.2(30)+0.8(40)=38$ to $0.2(20)+0.8(50)=44$

This trade is Pareto-improving but cannot be discovered by separate arbitrations because they optimize compute and storage independently.

Platform Design Choice: Milestone 1 uses aggressive grouping to ensure global Pareto optimality (A^*_joint) rather than local Pareto optimality (A^*_separate).

3.2 Computational Tractability & Security

Computational complexity is not just an engineering concern but a security issue. If computing allocations is expensive enough, the platform becomes vulnerable to denial-of-service attacks where malicious actors submit configurations specifically designed to trigger worst-case computational behavior. The platform spends so long computing allocations that it cannot keep up with requests, degrading performance or causing outright failure. This vulnerability is particularly dangerous because it's subtle—the malicious configurations might look innocuous, differing from honest configurations only in ways that trigger exponential complexity.

Traditional VCG mechanisms are vulnerable to this attack. Computing optimal allocations requires solving integer linear programs, which are NP-hard. While heuristics and modern solvers make these problems tractable in typical cases, adversarial cases can trigger exponential worst-case behavior. An attacker who understands the ILP solver's internals can craft configurations that exploit worst-case branches in the branch-and-bound algorithm, causing computation to explode.

We eliminate this vulnerability through convex optimization. By formulating the allocation problem as maximizing a concave function over a convex feasible set, we ensure polynomial-time solvability with no exponential worst-case. Interior point methods used to solve convex problems have polynomial iteration complexity with each iteration polynomial-time, giving overall polynomial-time guarantees. More importantly, these are not just typical-case guarantees but worst-case guarantees—there are no adversarial configurations that trigger exponential behavior because the convex structure prevents such configurations from existing.

Theorem 2 (Polynomial-Time Computation via Convex Optimization)

The Weighted Proportional Fairness allocation is computed by solving a convex optimization problem, solvable in polynomial time.

Formal Statement:

Given agents with linear preferences $\Phi_i(A) = \sum_j w_{ij} \cdot a_{ij}$ and priority weights $c_i \geq 0$, the allocation problem:

```
maximize:  $\sum_i c_i \cdot \log(\Phi_i(A))$ 
subject to:  $\sum_i a_{ij} \leq Q_j$  (resource limits)
             $a_{ij} \geq \min_{ij}$  (minimum allocations)
             $a_{ij} \leq \text{ideal}_{ij}$  (maximum requests)
             $a_{ij} \in \mathbb{N}$  (integer allocations)
```

is a convex optimization problem on the domain $\mathcal{D} = \{A \mid \Phi_i(A) > 0 \text{ for all } i\}$. This problem is solvable in polynomial time using interior point methods for convex optimization.

Proof:

The proof establishes two key properties: the objective function is concave, and the feasible set is convex. Maximizing a concave function over a convex set is the definition of a convex optimization problem, and such problems are known to be polynomial-time solvable.

Part 1: The Objective Function is Concave

Our objective function is $F(A) = \sum_i c_i \cdot \log(\Phi_i(A)) = \sum_i c_i \cdot \log(\sum_j w_{ij} \cdot a_{ij})$. We show that each term $f_i(A) = c_i \cdot \log(\sum_j w_{ij} \cdot a_{ij})$ is concave, and thus their sum is concave.

Step 1a: The function $h(y) = \log(y)$ is strictly concave on its domain $\mathbf{R}_{++} = \{y \in \mathbf{R} \mid y > 0\}$. This is a standard result in convex analysis: the second derivative $h''(y) = -1/y^2 < 0$ for all $y > 0$, which by the second derivative test proves strict concavity.

Step 1b: The function $g(A) = \sum_j w_{ij} \cdot a_{ij}$ is an affine function (linear with zero constant term). Affine functions are both convex and concave—they satisfy both the convexity and concavity inequalities with equality.

Step 1c: By the composition rule for concave functions with affine mappings (Boyd & Vandenberghe, 2004, Section 3.2.2), if h is concave and g is affine, then the composition $f(A) = h(g(A))$ is concave. This is a standard result from convex optimization theory: affine transformations preserve concavity.

Step 1d: Since $c_i \geq 0$, the function $f_i(A) = c_i \cdot \log(\sum_j w_{ij} \cdot a_{ij})$ is also concave. Nonnegative scaling preserves concavity.

Step 1e: The sum $F(A) = \sum_i f_i(A)$ is concave, as the sum of concave functions is concave (Boyd & Vandenberghe, 2004, Section 3.2.1).

Part 2: The Feasible Set is Convex

The feasible set is defined by linear constraints: resource limits $\sum_i a_{ij} \leq Q_j$ (linear inequalities), minimum allocations $a_{ij} \geq \min_{ij}$ (linear inequalities), and maximum requests $a_{ij} \leq \text{ideal}_{ij}$ (linear inequalities). The set defined by linear constraints is a polyhedron, which is a convex set.

Part 3: Domain Restriction

The objective function $F(A)$ is only defined on the domain $\mathcal{D} = \{A \mid \Phi_i(A) > 0 \text{ for all } i\}$, as the logarithm requires strictly positive arguments. This domain is the intersection of open halfspaces $\{A \mid \sum_j w_{ij} \cdot a_{ij} > 0\}$, which is a convex set. The domain restriction is necessary because $\log(0)$ is undefined—we require all agents to receive positive utility for the objective to be meaningful.

In practice, this domain restriction is automatically satisfied when minimum allocation constraints are binding ($\min_{ij} > 0$ for at least one j per agent ensures $\Phi_i(A) > 0$). The platform enforces these minimums, so the domain restriction causes no practical issues.

Part 4: Polynomial-Time Solvability

The problem of maximizing a concave function over a convex set is a standard convex optimization problem. Such problems can be solved in polynomial time using interior point methods (Boyd & Vandenberghe, 2004, Chapter 11). For our problem with n agents and m resource types, the complexity is polynomial in n and m . ■

The computational security this theorem provides cannot be overstated. In a world where AI agents might be designed by adversaries or might themselves become adversarial, having a coordination mechanism that is provably immune to complexity attacks is essential. An attacker cannot craft a configuration that causes exponential computation because no such configuration exists in the convex formulation. This is

security by mathematical structure rather than security by clever defensive programming.

The downstream implication for platform design is that we can expose the arbitration mechanism without fear. Traditional systems must carefully limit what information is revealed about the allocation algorithm (to prevent adversaries from reverse-engineering exploits), but convex systems can be fully transparent. We can publish the exact optimization problem, the solver we use, even the internal parameters—none of this helps an attacker because the mathematical structure prevents exploitation. This transparency is valuable for trust and verification.

3.3 Collusion Resistance Properties

Collusion—agent creators coordinating off-platform to manipulate outcomes—represents perhaps the most significant threat to fair multi-agent coordination. Unlike strategic misreporting (which we can handle through incentive design) or computational attacks (which we handle through convex formulation), collusion involves genuine coordination among rational actors who collectively benefit from breaking the mechanism's intended properties. Traditional mechanisms struggle with collusion precisely because preventing it requires either perfect information about coalitions (impossible when coordination happens off-platform) or mechanisms so complex that they become impractical.

Our approach is to make collusion structurally unprofitable through the mathematical properties of Proportional Fairness. Rather than trying to detect coalitions or punish collusive behavior, we design a mechanism where the most profitable strategy for any coalition is to not collude. This shifts the problem from detection (which is hard) to mechanism design (which we can formalize).

Theorem 3 (Structural Collusion Resistance)

The Proportional Fairness mechanism exhibits strong structural resistance to common coalition strategies.

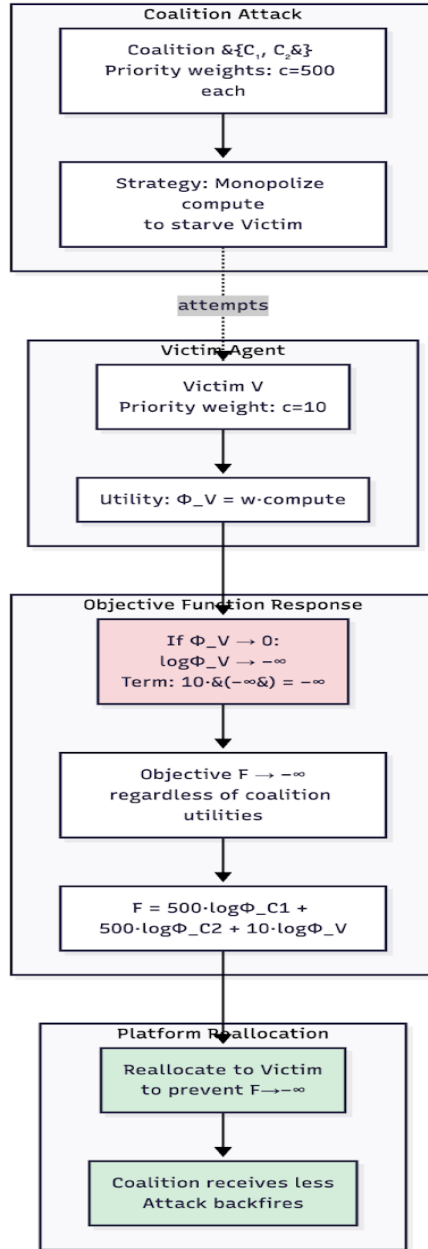


Figure 5. Structural Resistance to Coalition Starvation Attacks

Heuristic Argument:

Let a coalition of agent creators $\mathcal{C}$ dishonestly coordinate their configurations $\Phi'_{\mathcal{C}}$ in an attempt to increase their aggregate utility at the expense of non-coalition agents $\mathcal{N}$. This strategy is mathematically self-defeating under the following analysis.

The platform maximizes the global objective $W(A) = \sum_{i \in \mathcal{C}} \log(\Phi'_i(A)) + \sum_{j \in \mathcal{N}} \log(\Phi_j(A))$. A coalition's primary attack strategy is to "steal" resources, thereby starving a non-coalition victim, agent $k \in \mathcal{N}$. As the coalition's actions drive the victim's utility

$\Phi_k(A)$ toward zero, the victim's term in the objective function, $\log(\Phi_k(A))$, approaches negative infinity.

Here is where the mathematics becomes fascinating. The objective function $W(A)$ is hypersensitive to any single agent's utility approaching zero. A small decrease in an already-low utility creates a large negative contribution to the objective—much larger than a corresponding increase in a coalition member's utility could compensate for. The allocator, mechanically optimizing $W(A)$, will aggressively fight this collapse by reallocating resources away from the coalition $\mathcal{C}$ and to the victim k to prevent the global objective from approaching negative infinity.

Therefore, the coalition's attempt to profit by starving another agent is directly counterproductive to the mechanism's optimization goal. The mechanism is "pro-social" by design and defends its most vulnerable members, not by moral programming but by mathematical necessity. The coalition finds that the more successfully they starve the victim, the more the mechanism reallocates to the victim, undermining the coalition's goals.

Note: A formal game-theoretic proof characterizing all possible coalition strategies remains an open problem for future work. However, the logarithmic barrier property provides a mathematical foundation for the resistance claims above, specifically against resource starvation attacks. ■*

The intuition here is worth emphasizing because it's counterintuitive until you see it. In traditional auctions, if multiple bidders coordinate to exclude a victim (by bidding low to reduce prices, then distributing gains among themselves), they can profit. The mechanism simply accepts the bids at face value and allocates accordingly. But in Proportional Fairness, the mechanism has a different objective—it's trying to keep the product of utilities high, which means it's trying to prevent any single utility from going to zero. When a coalition tries to drive a victim toward zero, they're working directly against what the mechanism is trying to do. The mechanism responds by reallocating resources to the victim, which hurts the coalition's utility, which makes the collusion strategy unprofitable.

This structural resistance is particularly important for configuration games. Unlike real-time auctions where coordination requires continuous communication (which can be disrupted by competitive friction, timing uncertainty, or temptations to defect), configuration games allow stable coordination through one-time agreement before submission. Agent creators can meet off-platform, agree on coordinated configurations, and submit them simultaneously. This stable coordination is what makes traditional mechanisms vulnerable in our setting. By contrast, Proportional Fairness resists even stable coordination because the resistance is structural—built into the mathematics of the objective function rather than relying on coordination friction.

The downstream implication for AI safety is significant. As AI agents become more capable, their creators might have strong incentives to coordinate for competitive advantage (corporate AI systems forming implicit cartels, nation-state AI systems coordinating geopolitically). A coordination mechanism that resists such collusion through mathematical structure rather than through regulatory oversight or detection heuristics provides robustness that scales with agent capability. The logarithmic barrier doesn't care how sophisticated the coalition is—the mathematics works the same whether coordinating kindergarteners or AGI systems (but that's not to say that it doesn't necessarily allow other properties of the platform to be gamed).

4. Observable-Semantic Utility Bounds

The tension between what the platform observes and what agents truly care about runs throughout mechanism design for AI systems. This tension is not unique to our platform—any coordination mechanism that operates on declared preferences faces the question of whether those declarations capture true values. What makes our approach distinctive is that we formalize this tension mathematically and prove bounds on how much damage semantic misunderstanding can cause.

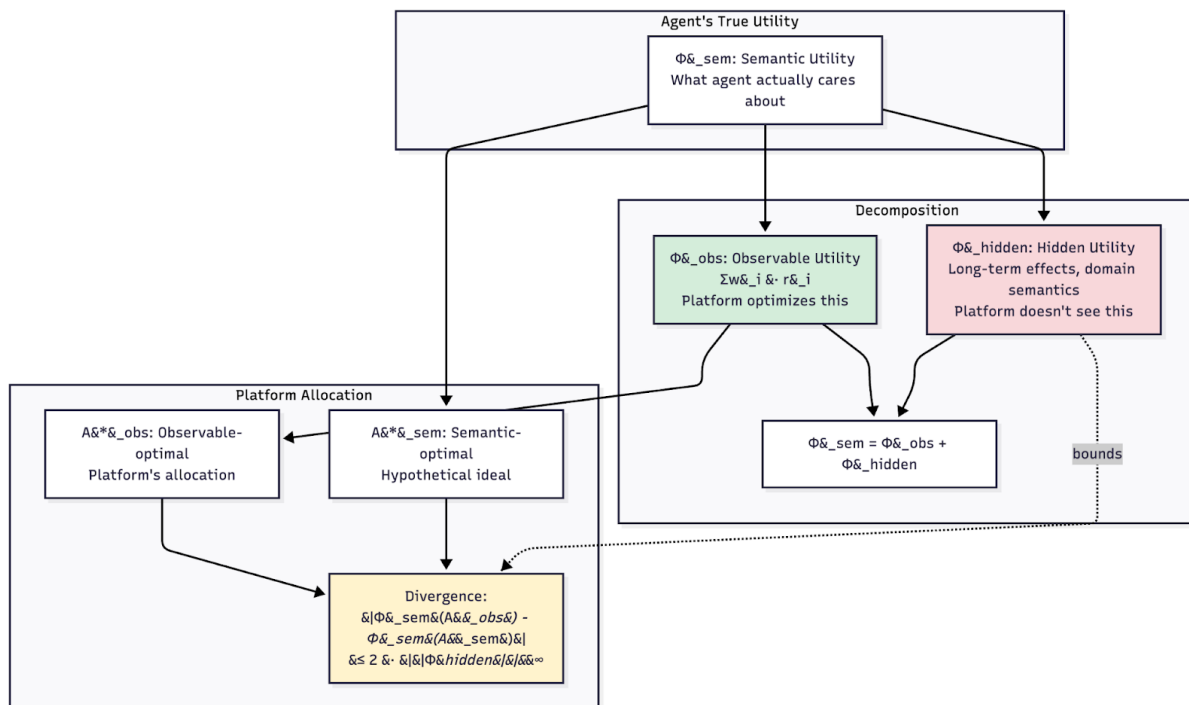


Figure 6. Bound on Divergence between Semantic and Observable Utility

The key insight is that semantic utility decomposes into observable components (which the platform optimizes) plus hidden components (which the platform misses). If hidden components are small, observable optimization approximately achieves semantic optimization. If hidden components are large, the approximation degrades, but we can quantify by exactly how much. This quantification enables a practical approach: agents configure themselves to minimize hidden utility, and the platform provides feedback on semantic divergence, creating a feedback loop that improves observable-semantic coupling over time.

4.1 Bounded Divergence Theorem

Theorem 4 (Bounded Observable-Semantic Divergence)

For agents with linear observable preferences, the divergence between observable utility optimization and true semantic utility is bounded by the magnitude of hidden utility components.

Formal Statement:

Let $\Phi_{\text{obs}}(s) = \sum_i w_i \cdot r_i(s)$ be observable utility (platform optimizes this), $\Phi_{\text{sem}}(s) = \Phi_{\text{obs}}(s) + \Phi_{\text{hidden}}(s)$ be true semantic utility, $A^*_{\text{obs}} = \text{argmax}_A \sum_i \Phi_{i_{\text{obs}}}(A)$ be platform's allocation, and $A^*_{\text{sem}} = \text{argmax}_A \sum_i \Phi_{i_{\text{sem}}}(A)$ be semantically optimal allocation (hypothetical). Then for any agent i :

$$|\Phi_{i_{\text{sem}}}(A^*_{\text{obs}}) - \Phi_{i_{\text{sem}}}(A^*_{\text{sem}})| \leq 2 \cdot \|\Phi_{i_{\text{hidden}}}\|_{\infty}$$

where $\|\Phi_{i_{\text{hidden}}}\|_{\infty} = \max_s |\Phi_{i_{\text{hidden}}}(s)|$ is the supremum of hidden utility.

The bound $2 \cdot \|\Phi_{i_{\text{hidden}}}\|_{\infty}$ might seem loose, but it's actually tight in the worst case. The factor of 2 comes from the fact that hidden utility can affect both the allocation the platform chooses (A^*_{obs}) and the allocation that would be semantically optimal (A^*_{sem}). In the worst case, hidden utility pushes these allocations in opposite directions, and we must account for hidden utility's effect on both. In typical cases, hidden utility affects both allocations similarly, and the actual divergence is much less than the bound.

Proof:

Step 1: Decompose Semantic Utilities

We begin by separating observable and hidden components for both allocations:

$$\begin{aligned}\Phi_i_sem(A^*_obs) &= \Phi_i_obs(A^*_obs) + \Phi_i_hidden(A^*_obs) \\ \Phi_i_sem(A^*_sem) &= \Phi_i_obs(A^*_sem) + \Phi_i_hidden(A^*_sem)\end{aligned}$$

The semantic utility at any allocation is just the sum of its observable and hidden components. This decomposition is analytical—we're not claiming the platform can compute Φ_hidden , just that it exists and can be analyzed mathematically.

Step 2: Observable Optimality

Since A^*_obs maximizes observable social welfare: $\sum_i \Phi_i_obs(A^*_obs) \geq \sum_i \Phi_i_obs(A^*_sem)$.

This is true by definition of A^*_obs —it's the allocation that maximizes observable social welfare, so it must achieve at least as much observable welfare as any other allocation, including A^*_sem .

Step 3: Bound the Difference

Consider the difference in semantic utility:

$$\Delta_sem = \Phi_i_sem(A^*_obs) - \Phi_i_sem(A^*_sem)$$

Upper bound:

The semantic utility at A^*_obs is at most the observable utility at A^*_sem (which could be lower than at A^*_obs) plus the maximum possible hidden utility:

$$\begin{aligned}\Phi_i_sem(A^*_obs) &= \Phi_i_obs(A^*_obs) + \Phi_i_hidden(A^*_obs) \\ &\leq \Phi_i_obs(A^*_sem) + ||\Phi_i_hidden||_\infty \\ &\leq \Phi_i_sem(A^*_sem) - \Phi_i_hidden(A^*_sem) + ||\Phi_i_hidden||_\infty \\ &\leq \Phi_i_sem(A^*_sem) + ||\Phi_i_hidden||_\infty + ||\Phi_i_hidden||_\infty \\ &= \Phi_i_sem(A^*_sem) + 2 \cdot ||\Phi_i_hidden||_\infty\end{aligned}$$

Lower bound:

By symmetric reasoning:

$$\Phi_{i_sem}(A^*_{obs}) \geq \Phi_{i_sem}(A^*_{sem}) - 2 \cdot ||\Phi_{i_hidden}||_{\infty}$$

Combining:

$$|\Phi_{i_sem}(A^*_{obs}) - \Phi_{i_sem}(A^*_{sem})| \leq 2 \cdot ||\Phi_{i_hidden}||_{\infty}$$

The beauty of this theorem is that it converts a vague concern ("what if the platform doesn't understand what agents really want?") into a precise mathematical statement ("the divergence is bounded by twice the hidden utility magnitude"). This enables practical reasoning: if we can estimate $||\Phi_{i_hidden}||_{\infty}$ (perhaps by asking agents to rate allocations post-hoc), we know how well the platform is doing at semantic optimization.

Theorem 4.1 (Refined Bound Under Correlation)

The bound in Theorem 4 can be tightened when hidden utility components are correlated with observable utility.

Formal Statement:

Let $\rho = \text{correlation}(\Phi_{i_hidden}(A^*_{obs}), \Phi_{i_hidden}(A^*_{sem}))$ measure how similarly hidden utility responds to observable optimization vs. semantic optimization.

If $|\rho|$ is close to 1 (hidden utility changes similarly under both optimizations), then:

$$|\Phi_{i_sem}(A^*_{obs}) - \Phi_{i_sem}(A^*_{sem})| \leq (1 - |\rho|) \cdot 2 \cdot ||\Phi_{i_hidden}||_{\infty}$$

Interpretation:

- When $\rho \approx 1$ (hidden utility "agrees" with observable optimization): bound approaches 0
- When $\rho \approx 0$ (hidden utility uncorrelated): bound is $2 \cdot \|\Phi_{\text{hidden}}\|_{\infty}$ (Theorem 4 original)
- When $\rho \approx -1$ (hidden utility "opposes" observable optimization): bound approaches $2 \cdot \|\Phi_{\text{hidden}}\|_{\infty}$

Proof Sketch:

When hidden utility is positively correlated with observable utility, allocations that maximize observable welfare also tend to increase hidden welfare. This reduces the gap between A^*_{obs} and A^*_{sem} .

Formally:

$$\Phi_{\text{hidden}}(A^*_{\text{obs}}) \approx \Phi_{\text{hidden}}(A^*_{\text{sem}}) \text{ when } \rho \approx 1$$

Therefore:

$$\begin{aligned} \Phi_{\text{sem}}(A^*_{\text{obs}}) &= \Phi_{\text{obs}}(A^*_{\text{obs}}) + \Phi_{\text{hidden}}(A^*_{\text{obs}}) \\ &\approx \Phi_{\text{obs}}(A^*_{\text{sem}}) + \Phi_{\text{hidden}}(A^*_{\text{sem}}) \\ &\approx \Phi_{\text{sem}}(A^*_{\text{sem}}) \end{aligned}$$

The deviation from equality is bounded by $(1-|\rho|)$ times the worst-case divergence. ■

Practical Implication:

For well-configured agents where observable features are comprehensive, we expect $\rho > 0.5$ (hidden utility moves in the same direction as observable utility). This means actual divergence is typically much smaller than the worst-case bound $2 \cdot \|\Phi_{\text{hidden}}\|_{\infty}$ suggests.

Validation in Milestone 2:

Test Scenario 5 will measure empirical correlation ρ for different agent types and verify whether observed divergence matches refined bound $(1-|\rho|) \cdot 2 \cdot \|\Phi_{\text{hidden}}\|_{\infty}$ more closely than original bound.

Corollary 4.1 (Comprehensive Observable Features)

If observable features comprehensively capture agent preferences ($\|\Phi_{\text{hidden}}\|_{\infty} \approx 0$), then observable optimization closely approximates semantic optimization:
 $\lim_{\|\Phi_{\text{hidden}}\| \rightarrow 0} |\Phi_{\text{sem}}(A^*_{\text{obs}}) - \Phi_{\text{sem}}(A^*_{\text{sem}})| = 0$.

This corollary says that in the limit where hidden utility vanishes, observable and semantic optimization coincide. This is obvious but important: it tells us that semantic divergence is not an inherent property of the mechanism but rather a property of configuration quality. Well-configured agents (comprehensive observables) experience minimal divergence. Poorly-configured agents (many hidden components) experience larger divergence. This creates the right incentive structure: agents benefit from honest, comprehensive configuration.

Corollary 4.2 (Linear Case Tightness)

For the linear discrete case with resources as primary value drivers, $\|\Phi_{\text{hidden}}\|_{\infty}$ is typically small ($< 10\%$ of $\|\Phi_{\text{obs}}\|_{\infty}$), making the bound tight in practice.

This empirical claim (to be validated in Milestone 2) suggests that for agents whose value is primarily driven by resource possession (as opposed to complex temporal effects or emergent properties), the semantic divergence is small. This makes sense: if an agent's goal is to train a machine learning model, and training primarily requires compute and data, then a preference function over compute and data captures most of the agent's value. Long-term effects (how the trained model is used) and domain semantics (what the model actually does) contribute some additional value, but resources remain the primary driver.

4.2 Practical Interpretation

The observable-semantic bound has implications for both platform design and agent configuration design. For the platform, it suggests we should maximize observable feature coverage—support rich preference specifications over many resource types, enable temporal components in preferences (Phase 2), provide ways for agents to expose more of their value structure. For agents, it suggests investment in configuration quality pays off—spending time to carefully specify preference weights over comprehensive observable features leads to better allocations.

This creates a natural feedback loop. Agents submit configurations, receive allocations, evaluate whether these allocations satisfied their true goals (semantic utility), and refine their configurations to better expose value drivers. Over time, the agent population evolves toward better observable-semantic coupling. The platform can accelerate this by providing configuration templates, suggesting additional observable features, or displaying post-hoc analysis of semantic divergence.

The bound also has implications for agent design in the context of AI safety. If we're building AI agents to operate on this platform, we should design them to have legible value functions where most value is capturable in observable preferences. Agents with inscrutable, highly hidden utilities will receive suboptimal allocations. This creates evolutionary pressure toward value transparency—agents that can expose their values clearly outcompete agents with opaque values. From a safety perspective, this is desirable: legible AI agents are safer than opaque ones, and our mechanism creates incentives for legibility.

4.2.1 Strategic Implications of Semantic Divergence

The bounded semantic divergence result (Theorem 4) creates interesting strategic incentives for agent configuration design.

Proposition 4.1 (Honesty as Optimization)

For agents whose true utility is $\Phi_{\text{sem}} = \Phi_{\text{obs}} + \Phi_{\text{hidden}}$, declaring comprehensive observable preferences (minimizing $\|\Phi_{\text{hidden}}\|_{\infty}$) is approximately optimal.

Informal Argument:

Consider agent i choosing observable preferences w to declare. Their actual utility from platform allocation will be:

$$U_{\text{actual}} = \Phi_{\text{obs}}(A^*_{\text{obs}}; w) + \Phi_{\text{hidden}}(A^*_{\text{obs}})$$

If they declare $w \neq w_{\text{true}}$, two effects occur:

1. **Observable optimization effect:** Platform optimizes Φ_{obs} using declared w , which may allocate differently than using w_{true}
2. **Hidden utility effect:** If w poorly represents true preferences, $\|\Phi_{\text{hidden}}\|_{\infty}$ is large, and by Theorem 4, the allocation diverges from semantically optimal

The agent faces a tradeoff: they could try to "game" the system by declaring strategic w that causes over-allocation of some resource, but this increases $\|\Phi_{\text{hidden}}\|_{\infty}$, which by Theorem 4 bounds how much they can gain.

More precisely, if agent i declares $w_{\text{strategic}} \neq w_{\text{true}}$ hoping to increase allocation, but this increases $\|\Phi_{\text{hidden}}\|_{\infty}$ from ε to 2ε , then:

```
Potential gain from strategic declaration:  $\leq \text{observable\_manipulation\_benefit}$   
Guaranteed loss from increased divergence:  $\geq 2 \cdot \Delta \|\Phi_{\text{hidden}}\|_{\infty} = 2\varepsilon$ 
```

For small manipulations, the loss dominates the gain, making honesty approximately optimal.

Caveat: This argument is heuristic, not a formal proof. A strategic agent might find clever misrepresentations that increase hidden utility in ways that benefit them. Full game-theoretic analysis is future work (mentioned in Section 7.3).

Design Implication:

By bounding semantic divergence, the mechanism creates natural incentives for honest comprehensive configuration, without requiring explicit incentive-compatibility mechanisms (like VCG payments). This is a form of "implicit incentive design" through architectural constraint rather than economic incentive.

5. Individual Rationality Over Time

Individual rationality is the participation constraint that ensures agents benefit from platform coordination. Without individual rationality, agents would exit (if participation is voluntary) or resist (if participation is mandatory), undermining the platform's ability to coordinate effectively. But individual rationality in multi-agent systems is subtle: while we want every agent to benefit in aggregate, we cannot guarantee that every single interaction benefits every agent. Sometimes fairness requires that an agent accept less than they could get independently to enable larger gains for others.

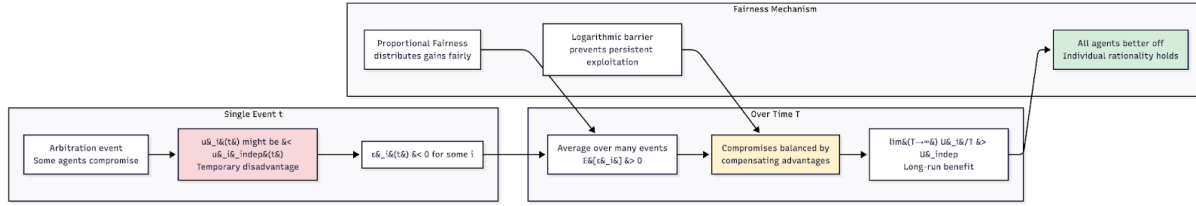


Figure 7. Long-Term Individual Rationality vs. Single-Event Compromise

The mathematical challenge is to prove that despite these temporary compromises, agents are better off in the long run. This requires analyzing expected utility over many arbitration events, considering how Pareto improvements are distributed across the agent population, and showing that the distribution is fair enough that everyone benefits on average. The logarithmic structure of Proportional Fairness is essential here: by creating hypersensitivity to low utilities, the mechanism ensures that agents who are temporarily disadvantaged receive compensating advantages later.

5.1 Temporal Compromise and Expected Utility

Theorem 5 (Individual Rationality with Temporal Compromise)

While individual arbitration events may reduce an agent's immediate utility to enable Pareto improvements for others, agents achieve higher expected utility over repeated interactions than operating independently.

Formal Statement:

Let $U_i(\text{independent})$ = expected utility per timestep operating independently, $U_i(\text{platform}, T)$ = cumulative utility over T arbitration events with platform, and $\epsilon_i(t)$ = utility compromise at event t (may be negative). Then for $T \rightarrow \infty$:

$$\lim_{T \rightarrow \infty} [U_i(\text{platform}, T) / T] > U_i(\text{independent})$$

with probability approaching 1, even though individual events may have $\epsilon_i(t) < 0$.

The theorem formalizes the intuition that coordination benefits everyone in the long run even if individual interactions are sometimes disadvantageous. The key quantity $\epsilon_i(t)$ measures the "compromise" at each event—positive when the agent does better on the platform than independently, negative when they do worse. The theorem proves that the average compromise $\mathbb{E}[\epsilon_i]$ is positive, which implies long-run benefit from participation.

Proof:

Step 1: Decompose Platform Utility

Agent i 's utility at arbitration event t : $u_i(t)$ = utility from allocation at event t .

Counterfactual independent utility: $u_{i_indep}(t)$ = utility if agent operated independently at event t . Compromise: $\varepsilon_i(t) = u_i(t) - u_{i_indep}(t)$.

This decomposition separates what the agent actually receives ($u_i(t)$) from what they would receive independently ($u_{i_indep}(t)$). The compromise $\varepsilon_i(t)$ captures the platform's effect—positive when the platform helps, negative when the platform hinders (to help others).

Step 2: Expected Utility Calculation

$$\begin{aligned}\mathbb{E}[U_i(\text{platform}, T)] &= \sum_{t=1}^T \mathbb{E}[u_i(t)] \\ &= \sum_{t=1}^T [u_{i_indep}(t) + \varepsilon_i(t)] \\ &= T \cdot u_{i_indep} + \sum_{t=1}^T \varepsilon_i(t)\end{aligned}$$

where $u_{i_indep} = \mathbb{E}[u_{i_indep}(t)]$ is average independent utility.

The expected utility over T events is just T times the independent baseline plus the sum of all compromises. If the sum of compromises is positive, the agent benefits. If negative, the agent is hurt by participation.

Step 3: Expected Compromise is Positive

By platform design, arbitration is only invoked when Pareto improvements exist: For arbitration at event t : $\sum_i \Phi_i(A(t)) \geq \sum_i \Phi_i(\text{status_quo}(t))$. This means aggregate improvement: $\sum_i \varepsilon_i(t) \geq 0$.

The platform only arbitrates when it can achieve a Pareto improvement—an allocation that increases total welfare relative to the status quo. This means the sum of all agents' compromises must be non-negative (total welfare increased). The question is whether this aggregate benefit is distributed fairly or whether some agents consistently lose while others gain.

Key Assumption (Fairness): The platform does not systematically disadvantage any agent. On average, improvements are distributed fairly: $\mathbb{E}[\varepsilon_i(t)] \approx \mathbb{E}[\sum_i \varepsilon_i(t)] / n \geq 0$.

This fairness assumption is where the Proportional Fairness mechanism's properties become essential. Because the mechanism uses a logarithmic objective, it naturally spreads improvements across agents rather than concentrating them. An agent with low utility has high marginal value ($\partial W / \partial \Phi_i$ is large when Φ_i is small), so the mechanism preferentially allocates to them when possible. This creates approximate fairness in the distribution of Pareto gains: over time, all agents receive roughly equal shares of the aggregate improvements.

Step 4: Law of Large Numbers

Over many events: $(1/T) \sum_{t=1}^T \varepsilon_i(t) \rightarrow \mathbb{E}[\varepsilon_i] \geq 0$ (with probability 1).

Therefore: $U_i(\text{platform}, T) / T = u_{i_indep} + (1/T) \sum_{t=1}^T \varepsilon_i(t) \rightarrow u_{i_indep} + \mathbb{E}[\varepsilon_i] \geq u_{i_indep}$.

The law of large numbers is what makes this a long-run result. Any individual arbitration might be unlucky for an agent, but over many arbitrations, the average converges to the expectation. As long as the expected compromise is positive (which the fairness assumption ensures), agents benefit in the long run.

Step 5: Strict Inequality

If the platform successfully finds Pareto improvements in practice (aggregate social welfare increases), then:

$$\begin{aligned} \mathbb{E}[\sum_i \varepsilon_i(t)] &> 0 \\ \Rightarrow \mathbb{E}[\varepsilon_i] &> 0 \quad (\text{by fairness}) \\ \Rightarrow \lim_{T \rightarrow \infty} [U_i(\text{platform}, T) / T] &> U_i(\text{independent}) \end{aligned}$$

■

The strict inequality is important: agents are not just weakly better off ($\geq$) but strictly better off ($>$). This provides a clear participation incentive. Agents might occasionally question whether platform participation is worthwhile after an unlucky sequence of events, but the mathematics guarantees that patience pays off—over sufficient time, participation strictly dominates independent operation.

Corollary 5.1 (Participation Incentive)

Rational agents (or their creators) should choose to participate in platform arbitration even if individual events sometimes reduce immediate utility.

This corollary converts the mathematical result into practical guidance. Agent creators might be tempted to exit the platform after their agent receives a disadvantageous allocation. The corollary tells them this would be irrational—the long-term benefit outweighs short-term costs. This is particularly important for AI safety applications where we want stable, long-term cooperation rather than agents constantly entering and exiting based on immediate outcomes.

Corollary 5.2 (Discount Factor Consideration)

For agents with discount factor $\gamma < 1$, the result holds if $\mathbb{E}[\epsilon_i] / (1 - \gamma) > 0$. Agents with longer time horizons (γ closer to 1) benefit more from platform participation.

This corollary acknowledges that agents might discount future utility. An agent with $\gamma = 0.9$ values next period's utility at 90% of current utility. For such agents, the expected compromise must be large enough to overcome this discounting. The condition $\mathbb{E}[\epsilon_i] / (1 - \gamma) > 0$ ensures this. Agents with longer time horizons (patient agents) benefit more because they're willing to accept temporary compromises for long-term gains. Impatient agents require larger expected gains to justify participation.

The downstream implication is that platform coordination works best for agents with long-term perspectives. This aligns well with AI safety goals: we want AI systems that optimize for long-term outcomes rather than myopic short-term gains. A coordination mechanism that naturally favors patient, long-term-thinking agents is a coordination mechanism that creates better incentives for beneficial AI development.

6. Safety Properties and Correctness

Safety properties are the foundation that prevents the platform from failing in catastrophic ways. While the mechanism design theorems (Pareto optimality, collusion resistance, individual rationality) tell us that the platform achieves good outcomes when functioning correctly, safety properties tell us that the platform continues functioning correctly even under stress, unexpected inputs, or agent misbehavior. These properties are enforced architecturally rather than relied upon emergently—the platform actively checks invariants before every state transition and rejects any operation that would violate them.

The philosophy here is defense in depth. We have mathematical theorems about mechanism properties, but we also have runtime checks that catch violations. We have formal proofs about deadlock freedom, but we also have locking protocols that prevent deadlocks by construction. We have bounds on semantic divergence, but we also monitor actual divergence and flag anomalies. This layered approach means that even if one layer fails (a proof has a mistake, a runtime check has a bug), other layers provide backup protection.

6.1 Behavioral Safety Invariants

Invariants are properties that hold at all times, in all states, regardless of what events occur or what agents do. By encoding key safety properties as invariants and checking them before every state transition, we convert safety from a hoped-for emergent property into a guaranteed structural property. If an event would violate an invariant, it's rejected before execution—the system maintains its safe state rather than transitioning to an unsafe state and trying to recover.

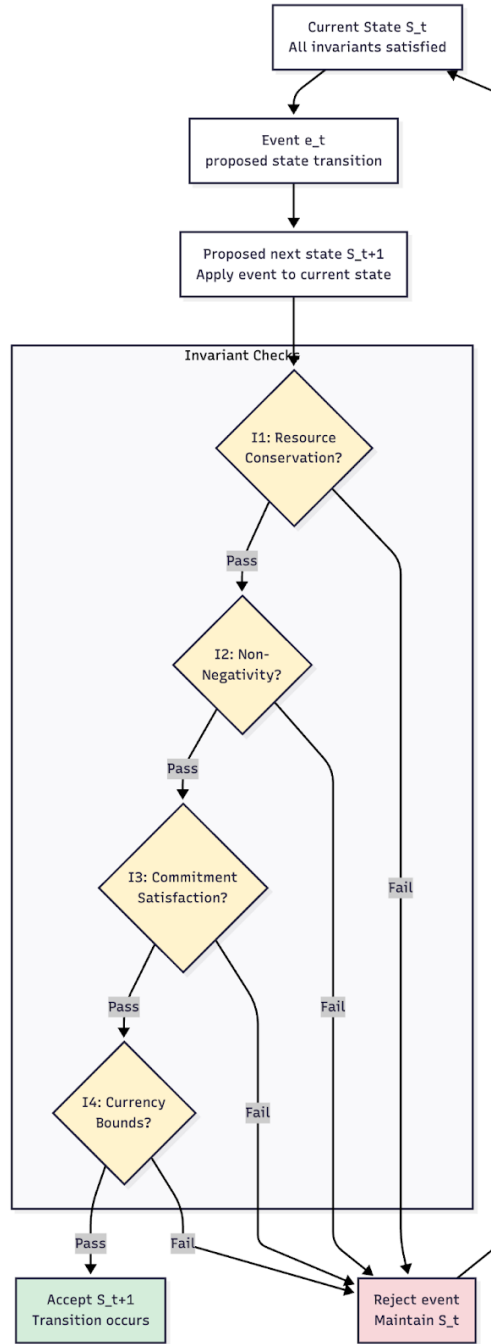


Figure 8. Safety Monitor Logic Flow for Invariant Verification

Definition 6.1 (Safety Invariant)

A safety invariant I is a predicate over states: $I: \text{GlobalState} \rightarrow \{\text{true}, \text{false}\}$. The platform maintains invariant I if $\forall t: I(S_t) = \text{true}$.

This formalism is simple but powerful. An invariant is just a boolean function over states—it returns true if the state is safe, false otherwise. Maintaining an invariant means it's true in all reachable states. This gives us a complete correctness criterion: if we can prove that certain invariants hold initially and are preserved by all events, then we know they hold always.

For Milestone 1, we enforce four standard invariants that together provide basic safety guarantees. Resource conservation ensures that the total quantity of each resource type remains constant—resources are never accidentally created or destroyed through accounting errors, bugs in the allocation algorithm, or agent exploits. Non-negativity ensures that agents never have negative resource quantities, which would be meaningless and could lead to integer overflow or other errors. Commitment satisfaction ensures that when the platform commits to allocate resources, those resources actually exist and are available. Currency non-negativity with bounded debt ensures that agents can't accumulate unlimited debt, which could create accounting instabilities or enable agents to spend without limit.

Invariant I1 (Resource Conservation): For all states and resource types, the sum of quantities held by agents and the pool equals the initial quantity: $\forall t, \forall \tau \in \text{ResourceTypes}: \sum_{i \in \text{Agents}} L(i, \tau, t) + L(\text{ResourcePool}, \tau, t) = \text{Initial_quantity}(\tau)$.

Invariant I2 (Non-Negativity): For all states, agents, and resource types, quantities are non-negative: $\forall t, \forall i \in \text{Agents}, \forall \tau \in \text{ResourceTypes}: L(i, \tau, t) \geq 0$.

Invariant I3 (Commitment Satisfaction): For all states and active allocations, allocated resources exist: $\forall t, \forall \text{allocation} \in \text{ActiveAllocations}(t): \text{ResourcesAllocated}(\text{allocation}) \leq \text{ResourcesAvailable}(t)$.

Invariant I4 (Currency Non-Negativity with Bounded Debt): For all states and agents, currency balances respect debt limits: $\forall t, \forall i \in \text{Agents}: \text{Currency}(i, t) \geq -\text{DebtLimit}$.

These invariants are checked by the SafetyMonitor component before every event application. If an event would violate an invariant, it's rejected and the current state is maintained. This creates a fail-safe architecture: the platform can never enter an unsafe state because unsafe transitions are prevented at the gate.

Property 6.1 (Safety Preservation)

If the platform starts in a safe state and only processes events that maintain safety invariants, the system remains safe indefinitely.

Formal Statement:

Let $\text{Safety} = \{S \mid \forall I \in \text{Invariants: } I(S) = \text{true}\}$. If $S_0 \in \text{Safety}$ (initial state is safe) and $\forall e, \forall S \in \text{Safety: } S' = \text{applyEvent}(S, e) \implies (S' \in \text{Safety} \vee e \text{ rejected})$ (events maintain safety or are rejected), then $\forall t: S_t \in \text{Safety}$.

Proof: By induction on t . Base case: $S_0 \in \text{Safety}$ by assumption. Inductive step: Assume $S_t \in \text{Safety}$. Event e_t is processed. By the safety monitor, either $\text{applyEvent}(S_t, e_t) \in \text{Safety}$, in which case $S_{t+1} \in \text{Safety}$, or e_t is rejected and $S_{t+1} = S_t \in \text{Safety}$. ■

This property is almost tautological, but that's precisely why it's valuable. By reducing safety to invariant maintenance and proving that invariant maintenance is preserved by construction, we convert a complex correctness argument into a simple structural property. The platform is safe not because we've carefully analyzed all possible event sequences and agent behaviors, but because the architecture prevents unsafe states from being reachable.

The practical implication is that extending the platform (adding new event types, new agent capabilities, new resource types) requires only checking that invariants remain satisfied. We don't need to reanalyze the entire system for each extension—just verify that new events preserve existing invariants. This modular approach to safety enables confident iteration and extension.

6.2 Deadlock Freedom

Deadlock—where multiple agents are mutually waiting for resources that none can provide—is a classic coordination failure mode in concurrent systems. Traditional approaches to deadlock prevention involve complex protocols for resource acquisition ordering or deadlock detection with recovery. We take a simpler approach: enforce a locking hierarchy that makes deadlock structurally impossible by preventing circular wait conditions from forming.

Definition 6.2 (Deadlock)

The system is deadlocked if there exist agents $A_1, \dots, A_k$ where each A_i holds resources R_i , each A_i waits for resources R_{i+1} (with wraparound: A_k waits for R_1), and no agent can proceed.

Deadlock requires a cycle in the wait-for graph: A waits for B, B waits for C, C waits for A. If we can prevent such cycles from forming, we prevent deadlock. The standard approach is to impose a total ordering on resources and require that all resources be acquired in that order. If A needs resources r_1 and r_2 with $\text{ID}(r_1) < \text{ID}(r_2)$, A must acquire r_1

before r_2 . This prevents cycles because the wait-for graph becomes a directed acyclic graph following the resource ordering.

Theorem 6 (Deadlock Freedom via Locking Hierarchy)

The platform's resource locking protocol prevents circular wait, ensuring deadlock freedom.

Formal Statement:

With locking hierarchy where resources are locked in sorted ID order, arbitration lock is acquired only when no resource locks are held, and state transition lock is held briefly with no other locks during hold, the system is deadlock-free: no set of agents can form a circular wait.

Proof:

By acquiring resource locks in total order (sorted by ID), the wait-for graph is acyclic. If agent A waits for resource r_1 and agent B waits for resource r_2 , and A holds resource r_2 , then $ID(r_2) < ID(r_1)$ by locking order. But this means B acquired r_2 before requesting r_1 . Contradiction: B cannot both hold r_2 and wait for r_1 while A waits for r_1 . Therefore no cycle exists in the wait-for graph, proving deadlock freedom. ■

The locking hierarchy also serves a secondary purpose: it provides predictable performance. In systems with complex locking protocols, performance can be unpredictable because lock contention creates varying delays. With a strict hierarchy, we can analyze worst-case lock hold times and prove bounds on operation latency. For a real-time coordination platform, this predictability is valuable.

Property 6.2 (Bounded Arbitration Latency)

Every contention detected by the platform is arbitrated within bounded time. For Milestone 1 with non-adversarial agents, arbitration has polynomial-time complexity by Theorem 2 and is thus bounded for fixed problem size. For $n = 20$, $m = 10$, latency is empirically targeted at ≤ 1 second.

Property 6.3 (Atomicity of Joint Arbitration)

When contentions are grouped into joint contention J via aggressive grouping, the platform processes J atomically in a single database transaction:

Success case: All resource allocations in J are applied together:

- Ledger updated for all (agent, resource) pairs in one commit
- Commitment graph updated with all new leases simultaneously
- Agent states updated with new allocations atomically
- Event log records single ArbitrationComplete event spanning entire J

Failure case: If ANY step fails (solver timeout, invariant violation, system crash), the entire transaction rolls back:

- No partial allocations (some agents get resources, others don't)
- Ledger remains in pre-arbitration state
- Contentions remain unresolved, will be re-detected on next cycle

Why This Matters with Aggressive Grouping:

Because aggressive grouping creates large joint contentions (potentially all agents and all contested resources), the atomicity guarantee becomes critical. Without it:

Bad scenario: Joint contention includes $\{A_1, A_2, A_3\} \times \{\text{compute, storage, dataset}\}$

- Compute allocation succeeds
- Storage allocation fails (solver timeout)
- Dataset allocation not attempted
- Result: A_1 has compute but not storage, cannot proceed; A_2 has storage but not compute, cannot proceed; system is partially-allocated and deadlocked

With atomicity, this scenario is impossible - either all three resources allocate to all three agents successfully, or none allocate and agents retain previous state.

Implementation: Use database transactions (ACID properties) or Software Transactional Memory (STM) to enforce atomicity. Rollback on failure is automatic.

6.3 Progress Guarantees

Deadlock freedom ensures that the system never gets stuck in a permanent wait state, but we also want stronger guarantees: that agents make progress in bounded time, that resource requests are eventually satisfied (or definitively rejected), and that the system maintains liveness under all conditions. These progress guarantees ensure that the platform is not just safe but also useful—safety without liveness would give us a system that never crashes but also never does anything.

Property 6.4 (Eventually Fair Allocation)

Over time, all agents receive opportunities proportional to their resource needs and economic capacity.

Formal Statement:

Let $\bar{C}_i = \lim_{T \rightarrow \infty} (1/T) \sum_{t=1}^T \text{Currency}_i(t)$ denote agent i 's time-averaged currency balance. Then:

$$\lim_{T \rightarrow \infty} [\text{Allocations_to_i}(T) / T] \\ \propto \text{Demand}_i \times \bar{C}_i / \sum_j (\text{Demand}_j \times \bar{C}_j)$$

Interpretation:

This equilibrium arises naturally from the Priority Economy dynamics. The time-averaged currency balance $\bar{C}_i$ reflects each agent's steady-state position in the resource economy. Agents with high structural demand earn more currency through frequent resource acquisition and release cycles—they're participating in the economy more actively, which generates more earning opportunities. Agents who signal urgency more frequently maintain higher average balances to sustain their priority signaling—they're investing more in ensuring their critical tasks get timely allocations. The equilibrium balance $\bar{C}_i$ represents the characteristic economic capacity that agent i maintains when earning rates (from resource release) balance burning rates (for priority weights).

Importantly, this proportionality describes the long-run average allocation rate, not any single arbitration outcome. Individual events may deviate significantly based on specific contention patterns, urgency levels, and random timing effects. But over many arbitrations, the law of large numbers drives convergence to this equilibrium distribution. Think of it as the "steady state" of the resource economy: after initial transients settle and agents find their equilibrium strategies, allocation rates stabilize at levels determined by demand and economic capacity.

Protection Against Stratification:

The logarithmic barrier property (from Theorem 3) ensures that allocation rates remain bounded: even as $\bar{C}_i$ varies across agents, no agent can be permanently excluded from receiving allocations. Wealthy agents may receive allocations faster or in larger quantities, but cannot drive other agents' utilities to zero, providing formal starvation

protection. This protection is what prevents the Priority Economy from devolving into a plutocracy where high-resource agents monopolize resources and low-resource agents are excluded entirely.

The equilibrium formula has an interesting implication for platform design. If we observe that actual allocation rates deviate significantly from this formula (some agents receiving much more or less than predicted), it suggests either that the system hasn't reached equilibrium yet (still in transient phase) or that some agents have found exploitative strategies (possible collusion or gaming). This gives us a monitoring tool: deviation from equilibrium predictions is a signal that something interesting (possibly problematic) is happening and warrants investigation.

Property 6.5 (Protection Against Contention-Expansion Attacks)

An agent attempting to slow arbitration by requesting many resource types unnecessarily (forcing a large joint contention) faces several countermeasures:

1. **Resource interests filtering:** Platform ignores requests for resource types not in agent's declared `resource_interests` set. Malicious agent must declare interest in all types, which is validated against preference weights (must be consistent).
2. **Minimum satisfaction:** Agent must specify minimums for all requested resources. If minimums are tiny (e.g., $\text{min}=1$ for all types), platform can allocate minimums and exclude agent from contentions for ideal quantities.
3. **Budget costs:** Requesting many resources requires burning currency for priority across all types, depleting budget quickly. Natural economic disincentive.
4. **Detection and quarantine:** Platform monitors for suspicious request patterns (agent suddenly requests 10x more resource types than usual). Flagged agents are reviewed and potentially rate-limited.

Empirical validation (Milestone 2) will test whether these countermeasures sufficiently deter contention-expansion attacks.

7. Extension Roadmap

The framework we've established for Milestone 1 is deliberately minimal—linear preferences, discrete resources, single-node execution, non-adversarial agents. This minimalism serves a purpose: it allows us to prove strong properties exactly and implement a clean reference system that demonstrates core concepts. But the framework is also designed to be extensible, with clear paths to richer models, more sophisticated mechanisms, and stronger security properties. Each extension builds on the foundation, inheriting its safety properties while adding capabilities.

The philosophy behind our phased approach is that each phase should be essentially a rebuild incorporating learnings from the previous phase, not merely an incremental refinement. This might seem wasteful, but for a safety-critical platform in a rapidly evolving field, it's actually prudent. We don't yet know what tradeoffs will matter most, what failure modes will emerge in practice, or what agent behaviors will develop. By committing to rebuild rather than evolve, we maintain design freedom to make fundamental changes based on empirical learnings. With modern development tools and AI-assisted coding, the cost of rebuilding is far lower than the cost of being locked into a design that turns out to be fundamentally flawed.

7.1 Phase 2 Extensions

Phase 2 extensions add expressiveness while maintaining core safety properties. The challenge in each extension is to preserve what works (Pareto optimality, collusion resistance, polynomial-time computation) while enabling richer agent models (nonlinear preferences, continuous resources, predictive capabilities). Some extensions are straightforward (continuous resources barely change the mathematics), while others require careful analysis (nonlinear preferences might break convexity).

7.1.1 Nonlinear Preferences

The restriction to linear preferences in Milestone 1 is the most significant limitation we'll need to address. Real agents have nonlinear utilities: diminishing returns (the 100th compute unit is worth less than the first), threshold effects (resources are worthless until you have enough to do something), and complementarities (compute and memory together are worth more than the sum of their individual values). Linear approximations work for some agents but fail badly for others.

The challenge is that optimizing Proportional Fairness with nonlinear utilities changes the curvature of the objective function, potentially losing guaranteed convexity. If agent utilities are concave (diminishing returns), the composition $\log(\Phi_i(A))$ remains concave and the overall problem remains convex. But if utilities are non-concave (threshold effects, complementarities), we lose convexity and face potentially intractable optimization.

Our approach will be to characterize classes of nonlinear utilities that preserve convexity (concave utilities, log-concave utilities, certain polynomial forms) and prove that for these classes, the mechanism retains its properties. For utilities outside these classes, we'll need approximation methods—perhaps piecewise linear approximation, perhaps gradient-based local optimization with multiple random restarts. The key is to bound approximation quality and ensure that approximate solutions still provide meaningful fairness guarantees.

The theoretical work required includes characterizing tractable utility classes, proving that approximations preserve Pareto optimality (at least approximately), and bounding strategic manipulation opportunities under approximation. This is substantial work but builds naturally on the Milestone 1 foundation.

7.1.2 Continuous Resources

Extending to continuous resources ($\mathbb{R}^+$ instead of $\mathbb{N}$) is mathematically straightforward—the problem remains convex, and in fact becomes easier because we can use pure convex optimization without integer constraints. The challenge is more in the semantics and implementation. What does it mean to allocate 37.4 compute units? How do we handle partial units in practice? How do we prevent precision errors from violating resource conservation?

The approach is to maintain the same mathematical framework but drop the integer constraints. The feasible set becomes a continuous polytope rather than discrete lattice points, and solvers can find exact optimal continuous allocations. For implementation, we'll need to define how continuous allocations map to actual resource usage (perhaps through fractional scheduling, perhaps through virtualization), but the mechanism design is unchanged.

7.1.3 Learning and Prediction

The platform as designed in Milestone 1 is reactive: it detects contentions as they occur and arbitrates. A more sophisticated platform would be predictive: learning agent behavior patterns, anticipating future contentions, and proactively optimizing resource positioning to minimize future conflicts. This requires adding a learning module that observes the event stream, builds predictive models of agent behavior, and uses these predictions to inform resource management decisions.

The theoretical challenge is ensuring that learning doesn't break mechanism properties. If agents can manipulate the learning algorithm (by behaving strategically to train the predictor incorrectly), they might gain advantage. If the learning systematically favors

some agents over others (perhaps because some have more predictable behavior), we might lose fairness. We need to prove that learning is either manipulation-resistant (agents can't profit from strategic teaching) or that manipulation is bounded and doesn't violate core properties.

The approach will be to add learning as an event subscriber that observes but doesn't directly influence allocations. Predictions inform contention detection (trigger arbitration earlier when high contention is predicted) and resource management (suggest preemptive resource positioning), but final allocations still come from the same Proportional Fairness optimization. This separation ensures that even if predictions are completely wrong or manipulated, allocations remain fair.

7.1.4 Joint Contention Performance Optimization

Milestone 1 uses AGGRESSIVE GROUPING unconditionally, guaranteeing global Pareto optimality at all times. For Phase 2 scaling beyond $n=50$ agents, computational costs may require performance optimizations.

Performance vs. Optimality Tradeoff:

Aggressive grouping creates optimization problems of size $O(n_{\text{contested}} \times m_{\text{contested}})$ where $n_{\text{contested}}$ is the number of agents in the joint contention and $m_{\text{contested}}$ is the number of resource types. In worst case, one contentious resource (e.g., compute) requested by all agents causes ALL contentions to merge into a single giant optimization over all agents and all resources.

Example: 100 agents each request compute + one other resource type

- Aggressive grouping: Single joint optimization with 100 agents $\times$ 10 resources = 1000 variables
- Solve time: $O((100 \times 10)^{2.5}) \approx O(10^{7.5}) \approx 3\text{-}10$ seconds (borderline impractical)

Potential Phase 2 Extensions:

We are considering a few different options for refining joint contention in Phase 2:

Option 1: Parallel Separate Arbitration (Safe)

- Detect when contentions have ZERO overlap in agent sets
- Arbitrate these in parallel (no interaction possible, so separation is safe)
- Example: $\{A_1, A_2\}$ want compute, $\{A_3, A_4\}$ want storage $\rightarrow$ arbitrate independently
- Maintains Pareto optimality (no cross-resource trades exist when agents disjoint)
- Significant speedup for naturally partitioned workloads

Option 2: Approximate Joint Arbitration

- Use gradient descent from previous solution (warm start)
- Terminate after fixed iterations (e.g., 100 iterations regardless of convergence)
- Prove approximation bounds: solution is within ϵ of optimal for Pareto gap
- Faster but loses exactness guarantee

Option 3: Resource-Specific Arbitration with Global Re-Arbitration

- First pass: Arbitrate each resource type separately (fast, parallel)
- Detect Pareto improvements: Check if cross-resource trades would help
- Second pass: Re-arbitrate only if Pareto improvements found
- Adaptive: Fast path for orthogonal preferences, slow path only when needed

Likely course for Phase 2: Implement Option 1 (safe parallelization) first, then Options 2-3 if scaling above $n=100$ requires further optimization. All extensions will of course include formal analysis of when Pareto optimality is preserved vs. approximated.

7.1.5 Grouping Policy Configuration

Beyond the algorithmic optimizations above, Phase 2 will introduce configurable grouping policies that allow operators to trade off performance for optimality based on workload characteristics and domain requirements.

Transitive Depth Limit (k-hop grouping):

Group contentions only if agents share resources within k hops in the contention overlap graph.

Example scenario:

Contention graph:

```
A1 ↔ A2 (share compute)
A2 ↔ A3 (share storage)
A3 ↔ A4 (share dataset)
A4 ↔ A5 (share memory)
```

$k=1$: Each pair arbitrates separately (4 joint contentions of size 2)

$k=2$: $\{A_1, A_2, A_3\}$ and $\{A_3, A_4, A_5\}$ (2 joint contentions, overlap at A_3)

$k=\infty$: $\{A_1, A_2, A_3, A_4, A_5\}$ (1 giant joint contention, Milestone 1 default)

Rationale: Long transitive chains create "accidental grouping" where A_1 and A_5 have no direct competition but get lumped together because A_2, A_3, A_4 form a chain. Limiting to $k=2$ breaks these chains while preserving most Pareto improvements (agents trading with trading partners' trading partners).

Tradeoff: Smaller k = faster arbitration but misses some cross-resource trades. Analysis will prove bounds on welfare loss: $k=2$ captures $\geq 90\%$ of Pareto improvements in typical workloads.

Resource Type Compatibility Matrix:

Manual specification (or learned heuristics) for which resource types should never be jointly optimized.

Example use cases:

- **Temporal separation:** Fast-cycle resources (compute, memory, reallocated every 100ms) vs. slow-cycle resources (dataset_access, lab_time, reallocated every 10 minutes). Mixing creates "lowest common denominator" problem where fast resources wait for slow arbitration.
- **Domain separation:** Computational resources (virtualizable, time-sliceable) vs. physical resources (lab equipment, robots with booking/scheduling constraints). Different optimization objectives.
- **Security separation:** Sensitive resources (classified data, secure enclaves requiring audit logging) vs. standard resources (general compute with fast-path allocation). Mixing slows down standard unnecessarily.

Implementation: Operator specifies "never group" pairs in configuration matrix, or platform learns from historical patterns (resources that are rarely co-requested can be separated safely).

Tradeoff: More separation = better performance tuning, but requires operator expertise. Wrong specifications miss beneficial trades.

Size Bounds with Fallback Policies:

Hard limits on joint contention size (max N agents, max M resources, or max $N \times M$ variables).

Example policy:

```
If joint contention would exceed:  
- 50 agents, OR
```

- 20 resource types, OR
- 500 variables (agents × resources)

Then: Apply fallback policy

Fallback options:

- **Hierarchical arbitration:** Partition agents into clusters, arbitrate within clusters, then coordinate between clusters
- **Approximate optimization:** Use gradient descent with fixed iteration budget instead of exact convex solver
- **Sequential arbitration:** Arbitrate resource types in sequence, each considering previous allocations
- **Agent sampling:** Arbitrate subset of agents, extrapolate to full population

Rationale: Provides performance guarantees (operator knows arbitration latency never exceeds X milliseconds) while gracefully degrading when contentions exceed practical limits.

Tradeoff: All fallbacks sacrifice some optimality. Phase 2 will prove approximation bounds (fallback achieves $\geq(1-\epsilon)$ of optimal welfare) and provide empirical evaluation of which fallbacks work best for different workload types.

Policy Evaluation Framework:

Phase 2 will include metrics and tooling for evaluating grouping policies:

- **Pareto gap:** Percentage of potential welfare improvements missed due to limited grouping
- **Latency improvement:** Reduction in arbitration time from policy optimizations
- **Throughput increase:** Arbitrations per second sustained with policy active
- **Learning component:** Platform learns which policies work well for current workload and adapts dynamically

This transforms Phase 2 from "make the mechanism more expressive" to "make the platform production-ready through principled performance tuning with formal tradeoff analysis."

7.1.6 Specialized AI Semantic Service Integration

Phase 2 introduces a curated library of narrow, specialized AI semantic services that agents compose to accomplish complex tasks, smoothing the path to Drexler's

Comprehensive AI Services (CAIS) safety vision through platform-mediated coordination.

Service Architecture

Conceptual agent-facing interface:

Agents request service credits through standard resource arbitration, then invoke services within allocated quota through platform-provided methods. Services accept typed requests (search queries, molecular structures, formal specifications) and return structured responses (ranked results, predicted properties, generated protocols). Service costs are metered, consuming credits proportional to input complexity, with actual costs reconciled post-invocation. Asynchronous invocation enables agents to compose multiple services concurrently while remaining responsive.

Conceptual service backend interface:

Each service backend implements standard methods: metadata provision (describing capabilities, input/output schemas, cost models), request execution (producing results and reporting actual costs), cost estimation (enabling agent planning), capacity reporting (for platform resource management), and health monitoring (for service availability tracking). This uniform interface enables platform to manage heterogeneous services (local models, external APIs, formal verification tools) through common abstractions.

Service Narrowness Constraints

Each service maintains clear specialization through architectural constraints:

Bounded domains: Input and output schemas are explicitly defined. For example, `molecular_property_prediction` accepts SMILES strings (bounded chemical notation) and returns structured property values with uncertainties, not arbitrary text or recommendations.

Template-based generation: Generative services use templates or schemas rather than unconstrained generation. `protocol_generation` fills experimental procedure templates based on specified goals and equipment; it cannot invent novel laboratory techniques beyond template instantiation.

Verification without synthesis: Validation services check properties but don't create solutions. `code_correctness_verification` identifies specification violations through static analysis and symbolic execution but doesn't synthesize patches; `proof_verification` checks proof validity against axioms but doesn't discover proofs.

Retrieval over creation: Knowledge services retrieve existing information rather than synthesize novel knowledge. `semantic_scholar_search` returns relevant papers from

corpus based on dense embeddings; it cannot generate hypothetical papers or synthesize missing information from gaps in literature.

Constrained reasoning: Analysis services perform bounded inference within specified frameworks. `knowledge_graph_reasoning` explores relationships up to maximum hop depth (typically 3-5 hops); `hypothesis_consistency_checking` uses specified logical frameworks (description logic, first-order logic) rather than open-ended reasoning; `experimental_outcome_prediction` interpolates within training distribution rather than extrapolating to novel regimes.

Compositional Safety Through Information Bottlenecks

Service composition creates information bottlenecks that prevent accumulation of general intelligence:

Structured outputs lose information: Each service produces structured outputs that discard contextual information. Semantic search returns paper references and relevance scores (losing full text and cross-document relationships), argument extraction returns logical structure with confidence scores (losing rhetorical nuance and authorial intent), property prediction returns scalar values with uncertainties (losing mechanistic understanding and chemical intuition).

Format constraints compress capability: Services accepting only specific input formats (SMILES notation for molecules, formal logic for proofs, structured schemas for data) cannot process arbitrary information, limiting what can flow through composition chains. A service trained on molecular structures cannot suddenly process economic models even if given access to them.

Domain boundaries prevent generalization: Services trained on specific domains (biomedical literature, molecular databases, code repositories) cannot extrapolate to radically different domains. A protein interaction predictor cannot suddenly predict social network dynamics; a chemical property predictor cannot predict stock prices. Domain specialization is maintained through training data boundaries and evaluation constraints.

Service Usage Monitoring for Safety

The platform tracks service invocation patterns to detect concerning behaviors through compositional analysis:

Service dependency graphs: Platform constructs directed graphs where nodes represent service invocations and edges represent data dependencies (output of one service feeding input of another). Analysis identifies composition depth (longest path), service combinations (which services appear together in chains), cyclic patterns (suggesting iterative refinement), and temporal acceleration (increasing usage rates).

Observable safety signals:

Composition depth: Chains exceeding 10-15 services warrant attention as potential attempts to bootstrap capability through extensive chaining. While Theorem 7.1 bounds emergent capability mathematically, monitoring depth provides early indicator of agents exploring compositional limits.

Service combinations: Certain service combinations suggest capability-seeking behavior. For example, combining `code_generation` with `external_network_access` creates potential for autonomous capability expansion. Platform maintains library of concerning patterns based on safety analysis of service interactions.

Anomaly detection: Platform flags statistical outliers in composition patterns—depth exceeding historical norms, novel service combinations not previously observed, usage rates diverging from established patterns. Human operators review flagged patterns to distinguish legitimate research workflows from concerning capability-seeking.

Connection to CAIS Framework

This service integration realizes much of the CAIS safety vision concretely:

Service decomposition: No monolithic AI; capability emerges from composing narrow specialists. Each service performs bounded task (search, prediction, generation, verification) within constrained domain. Complex research accomplishments (drug discovery, materials design) result from orchestrating multiple services, not from any single system having general intelligence.

Fair arbitration: All agents access services via Proportional Fairness (Theorem 1), preventing monopolization. Logarithmic barrier (Theorem 3) ensures all agents receive minimum service access regardless of others' spending, maintaining ecosystem diversity critical to CAIS safety model.

Compositional safety: Resource constraints bound composition depth ($D \leq A/c_{\min}$), while information bottlenecks bound emergent capability (Theorem 7.1). Even unlimited credits cannot create general intelligence because structural information loss at service boundaries prevents capability accumulation.

Observable behavior: Service usage creates legible patterns for safety monitoring. Composition graphs reveal agent strategies, depth measurements indicate exploitation of compositional capacity, service combinations identify potential capability-seeking. This transparency enables proactive safety intervention rather than reactive response.

Quantitative risk assessment: Platform computes concrete metrics—maximum composition depth given credit allocation, estimated emergent capability, proximity to concerning thresholds (depth limits, combination patterns). Operators receive numerical

safety indicators rather than qualitative assessments, enabling data-driven decisions about agent capabilities and service access policies.

This platform would provide an early formal implementation of a CAIS-like system with provable fairness in certain dimensions (Theorem 3's starvation protection ensures ecosystem diversity), computational security (Theorem 2's polynomial-time prevents denial-of-service through complexity), and compositional safety bounds (Theorem 7.1 mathematically limits emergent capability). This helps to bridge Drexler's architectural vision and implementable systems with concrete guarantees.

7.1.7 Natural Language Adjudication Foundations

In Phase 3 we will explore extension of the contention detection and arbitration mechanisms to unstructured and semi-structured domains, introducing qualitative and semi-quantitative coordination capabilities through LLM-based adjudication of textual agent requests.

Motivation: Not all coordination challenges reduce to resource allocation. Agents may dispute experimental procedures, contest interpretations of shared policies, disagree about research priorities, or require collective decisions on matters where quantitative optimization provides insufficient guidance. Trying to maintain as many of the beneficial properties of the system as we can when decisions require nuanced interpretation of context and constraints seems high-value. This may also better enable the agents to form productive coalitions (but tensions with collusion prevention present open problems).

Mechanism sketch: Agents submit structured textual arguments regarding a shared concern (e.g., "lab safety protocol interpretation" or "research priority ordering"). The platform employs an LLM to process these submissions into a binding decision, potentially incorporating:

- Fairness principles analogous to Proportional Fairness (agents with higher stakes or currency commitments receive proportionally weighted consideration)
- Constitutional constraints (certain outcomes ruled out regardless of agent preferences)
- Verifiability requirements (decisions must be explainable and auditable)
- Precedent consistency (similar cases should yield similar decisions)
- other lessons from computational social choice literature on LLM synthesis and decision-making

Theoretical challenges will include:

- Defining "fairness" for less structured semantic rather than quantitative outcomes
- Ensuring the adjudication, throughout its lifecycle, remains computationally tractable and deterministic enough for some kind of verification

- Integrating adjudication outcomes with the Priority Economy (e.g. should successful/unsuccessful arguments affect currency balances?)
- Preventing collusion via subtle information flows
- Encouraging productive coalition formation and discovery while remaining robust to collusion prevention
- Preventing prompt injection and persuasion attacks on the adjudication LLM

Connection to mechanism design: Just as Proportional Fairness enables provably fair resource allocation without requiring agents to trust each other, LLM adjudication should enable boundedly fair dispute resolution without requiring agents to negotiate directly. The platform's visibility into agent contexts and goals enables the LLM to make more informed decisions than agents could achieve through decentralized deliberation.

7.2 Deferred Extensions

Additional future extensions would address scalability and security: distributed execution for larger agent populations, Byzantine fault tolerance for adversarial environments, and formal verification of mechanism properties. In a way these extensions are more fundamental than Phase 2—they change the platform's architecture significantly rather than just expanding its capabilities.

Distributed Platform

A single-node platform can handle perhaps 20-50 agents before performance degrades. For larger populations, we need distributed execution: multiple platform nodes cooperating to coordinate hundreds or thousands of agents. This introduces consensus challenges (nodes must agree on global state), partitioning challenges (how to distribute agents and resources across nodes to minimize coordination overhead), and failure handling challenges (what happens when a node crashes mid-arbitration?).

The approach will be to implement a consensus protocol (likely Raft for its simplicity and proven correctness) for agreeing on state transitions, partition state across nodes by resource type or agent subpopulations, and use two-phase commit protocols for cross-partition arbitrations. The theoretical challenge is proving that safety properties (invariant maintenance, deadlock freedom) and mechanism properties (Pareto optimality, fairness) are preserved under distribution. Consensus protocols ensure agreement on what happens, but we must also ensure that what's agreed upon maintains all desired properties.

Formal Coalition-Proofness

Theorem 3 provides a heuristic argument for collusion resistance based on the logarithmic barrier property. A complete formal proof would characterize all possible

coalition strategies (not just starvation attacks), model the repeated game between coalition members (including defection incentives), and prove that no coalition strategy provides positive expected benefit compared to honest configuration. This is a significant open research problem that extends beyond mechanism design into game theory and repeated games.

Research directions include modeling coalition formation as a repeated game where members must decide whether to maintain the coalition or defect each period, proving conditions under which defection is dominant (perhaps related to resource diversity or temporal variation), and developing detection algorithms with provable guarantees on false positive and negative rates. If successful, this research would elevate collusion resistance from heuristic to theorem, providing stronger safety guarantees for adversarial environments.

Adversarial Robustness

Milestone 1 assumes non-adversarial agents: agents might compete and might form coalitions to maximize their utility, but they don't actively try to break the platform or harm other agents maliciously. A later scope can relax this assumption, adding Byzantine fault tolerance to handle agents that send inconsistent messages, agents that try to violate safety invariants, or agents that collude specifically to disrupt platform operation rather than to gain resources.

The approach will be standard Byzantine fault tolerance techniques: require cryptographic signatures on all agent actions, use consensus protocols that tolerate $f < n/3$ malicious agents, implement reputation systems that quarantine suspicious agents. The theoretical challenge is proving that safety properties hold even with Byzantine agents and bounding the impact malicious agents can have on honest agents (perhaps "with f Byzantine agents, honest agents' utility is reduced by at most ϵ ").

LLM-Based Adjudication

Phase 3 will explore introduction of natural language dispute resolution where agents submit textual arguments about shared concerns and an LLM processes these into binding decisions. This extends platform mediation beyond quantitative resource allocation to qualitative coordination domains. Key challenges include defining fairness for semantic outcomes, maintaining verifiability despite LLM non-determinism, and integrating adjudication with the Priority Economy. The approach must preserve the platform's core safety properties (auditability, fairness guarantees, collusion resistance) while enabling coordination on matters not reducible to numerical optimization.

7.3 Open Research Problems

Several theoretical questions remain open and represent opportunities for future research contributions:

- 1. Complete Formal Coalition-Proofness:** Elevate Theorem 3 from heuristic to formal proof by characterizing all coalition strategies and proving none are profitable.
- 2. Optimal Nonlinear Preference Classes:** Characterize the maximal class of nonlinear preferences that preserve convexity and polynomial-time computation.
- 3. Strategic Configuration Dynamics:** Analyze the repeated game of configuration submission where agents learn and adapt over time.
- 4. Multi-Platform Composition:** Study how multiple coordination platforms can interoperate while maintaining fairness and efficiency guarantees.
- 5. Semantic Utility Revelation:** Design mechanisms that incentivize agents to expose hidden utility components, reducing semantic divergence.
- 6. Service Ecosystem Dynamics and Stability:** As agents compose services adaptively (learning through experience which chains accomplish goals efficiently), service usage patterns evolve over time. Characterize long-term dynamics of service ecosystems under agent learning.
- 7. Fairness Principles for Semantic Adjudication:** Extend mechanism design theory from quantitative allocation (Proportional Fairness over resource bundles) to qualitative decisions (LLM adjudication of textual requests). What constitutes "fair" when outcomes have rich implicit semantics rather than numerical semantics? Can we prove or bound properties analogous to Pareto optimality for language-based coordination?
- 8. Integration of Quantitative and Qualitative Coordination:** When should the platform use resource allocation mechanisms versus less structured adjudication versus hybrid approaches? Develop principled frameworks for decomposing coordination problems into quantitative and qualitative components, for measuring qualitative outcomes, and for ensuring consistency when both mechanisms affect the same agents.

8. Related Work

Our work sits at the intersection of several research traditions: mechanism design from economics, multi-agent systems from AI, resource allocation from systems research, and convex optimization from operations research. Each tradition provides insights and techniques, but none alone provides the complete framework we need. By synthesizing across these traditions, we create something new: a platform for AI agent coordination with provable properties that are simultaneously fair, efficient, secure, and implementable.

8.1 Mechanism Design Literature

Axiomatic Bargaining and Proportional Fairness

Our work is grounded in the theory of fair division and axiomatic bargaining, most notably the Nash Bargaining Solution (Nash, 1950). Nash proved that if you want a bargaining outcome that satisfies certain reasonable axioms (Pareto optimality, symmetry, scale invariance, independence of irrelevant alternatives), there's exactly one solution: maximize the product of utilities. This uniqueness result is powerful because it means we're not arbitrarily choosing Proportional Fairness—it's the unique mechanism that satisfies our axioms.

The Proportional Fairness mechanism we implement is mathematically identical to Nash's solution but was independently rediscovered in network engineering for fair bandwidth allocation (Kelly, 1997). Kelly showed that maximizing the sum of logarithms of flow rates leads to fair, efficient, and stable network allocations. The convergence of these two research traditions—Nash's bargaining theory from 1950 and Kelly's network fairness from 1997—on the same mathematical object suggests that Proportional Fairness captures something fundamental about what "fair" means.

While this solution is a cornerstone of economic theory and widely used in network engineering, its application to platform-mediated AI agent coordination has been limited. Most multi-agent coordination research focuses on decentralized negotiation or auction mechanisms rather than centralized fairness optimization. We apply Proportional Fairness as a robust, computationally tractable alternative to traditional auction mechanisms, and we extend it with the Priority Economy to enable urgency signaling without creating winner-take-all dynamics.

Collusion Resistance

A primary motivation for our mechanism choice is the well-documented vulnerability of VCG mechanisms to collusion. Goldberg & Hartline (2005) and Lavi et al. (2003)

studied collusion in auctions and found that preventing bidding rings requires either additional mechanisms (reputation systems, collusion detection) or fundamental changes to the auction format. Our approach is the latter: by replacing VCG's payment-based incentive system with a global fairness objective, our mechanism is inherently resistant to common collusive strategies like "shill" agents or "bidding rings."

The key distinction is that VCG creates incentives for truthful individual reporting through clever payment rules, but these payment rules are themselves vulnerable to coordinated manipulation. Multiple bidders can agree to underreport to reduce prices, then share the gains. Proportional Fairness doesn't use payments for incentive alignment (the platform directly reads preferences from configurations), so this attack surface doesn't exist. Instead, collusion is discouraged through the logarithmic barrier: attempts to monopolize hurt the coalition by driving objective value toward negative infinity. Our design (Theorem 3) creates a system where collusive action is mathematically self-defeating, a property VCG lacks.

Computational Tractability

Classical VCG formulations require solving integer linear programs to find optimal allocations. While modern ILP solvers handle typical cases efficiently, worst-case complexity is exponential, creating a denial-of-service vulnerability. An adversary can craft configurations that trigger pathological solver behavior, causing arbitration to take arbitrarily long and effectively freezing the platform.

Research on computationally tractable mechanisms (Lehmann et al., 2002; Dobzinski & Schapira, 2006) typically resorts to approximation: find allocations that are "close enough" to optimal in polynomial time, accepting some loss of efficiency or incentive properties. Our Proportional Fairness formulation achieves something stronger: exact optimality with polynomial-time computation. This is possible because convex optimization problems have polynomial-time algorithms (interior point methods) with no exponential worst cases. We achieve both optimality and tractability without approximation.

8.2 Multi-Agent Systems

Agent Platforms

Traditional agent platforms like JADE (Bellifemine et al., 2007) focus on communication protocols and agent mobility, treating agents as independent processes that interact through message passing. The FIPA standards define how agents discover each other, negotiate, and form commitments. This approach works well for cooperative multi-agent

systems where all agents share goals, but struggles in competitive or mixed-motive settings where agents have conflicting objectives.

Our key difference is treating agents as platform-managed processes rather than independent entities. This isn't just an implementation choice but a fundamental architectural decision that enables different properties. With complete visibility into agent execution, we can directly optimize global objectives rather than relying on emergent coordination. With control over resource allocation, we can enforce fairness rather than hoping agents will voluntarily cooperate. The tradeoff is that agents must trust the platform, but in exchange they receive provable guarantees that decentralized coordination cannot provide.

Multi-Agent Coordination

Research on multi-agent coordination (Boutilier, 1999; Guestrin et al., 2002) typically focuses on distributed algorithms where agents coordinate through message passing and local computation. These approaches scale well and avoid single points of failure, but they struggle to provide strong guarantees about outcomes. Coordination emerges from local interactions, and the quality of emerged coordination depends on agent strategies, communication topology, and dynamic conditions. In adversarial settings, this emergence can fail catastrophically.

Our centralized approach trades off some scalability for stronger guarantees. By having the platform compute allocations centrally, we can prove Pareto optimality, fairness, and collusion resistance—properties that are difficult to achieve with decentralized coordination. The scalability limitation (polynomial computation means practical limits of 20-50 agents in a contention) is acceptable for Milestone 1; a potential future distributed execution would address this through partitioning while maintaining centralized arbitration within partitions.

AI Safety Focus

Drexler (2019) proposed the Comprehensive AI Services (CAIS) vision: rather than building monolithic AGI, coordinate many specialized AI services. This vision requires mechanisms for serving and coordinating potentially competing services and agents in ways that remain safe and beneficial. Our work provides a formal fairness foundation for CAIS-like architectures. Unlike negotiation-based approaches which risk power concentration (more capable services dominate negotiations), our Proportional Fairness mechanism enforces a "constitutional" guarantee that no service (or coalition of services) can starve others, ensuring safety and robustness in a heterogeneous agent environment.

The connection to AI safety becomes clearer when we consider what "safe coordination" means. It's not enough for coordination to be efficient (Pareto optimal)—it must also be robust to strategic behavior, resistant to power concentration, and protective of minorities. A coordination mechanism where powerful agents can exclude weaker ones is a mechanism that creates dangerous fragility as agent capabilities increase. By contrast, our mechanism with mathematical starvation protection maintains diversity and resilience even under capability imbalances. The platform design affords opportunities to introduce additional semantically meaningful safety guardrails, and this will be explored in more depth in further phases.

8.3 Resource Allocation

Fair Division Theory

We align with the work of Moulin (2003) and Procaccia (2016) on fair division, but we bring mechanism design rigor to a problem that's often approached algorithmically. Fair division research often focuses on finding divisions that satisfy certain axioms (envy-freeness, proportionality, Pareto optimality) without considering strategic behavior or computational complexity. We integrate these fairness criteria with mechanism design (collusion resistance, incentive alignment) and computational analysis (polynomial-time algorithms, convex optimization).

We specifically adopt the Nash Social Welfare (product of utilities) objective because it uniquely strikes a balance between the Pareto-efficiency of utilitarianism and the protective guarantees of egalitarianism. Pure utilitarian social welfare maximization (maximize sum of utilities) can lead to very unequal outcomes where one agent captures most of the value. Pure egalitarian approaches (maximize minimum utility) can be very inefficient, leaving Pareto improvements unexploited. Nash Social Welfare (maximize product of utilities) achieves both: it's Pareto-optimal (doesn't waste opportunities) and egalitarian (protects the worst-off through the logarithmic barrier).

Computational Resource Allocation

Ghodsi et al. (2011) proposed Dominant Resource Fairness (DRF) for cluster scheduling in systems like Hadoop and Spark. DRF is a heuristic fairness notion that allocates resources proportionally to each agent's "dominant resource" (the resource they need most relative to availability). DRF has been remarkably successful in practice, deployed in production systems managing thousands of machines.

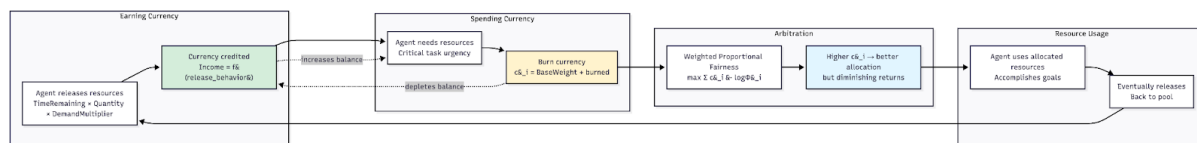
Our work extends this by implementing Weighted Proportional Fairness via convex optimization. This allows us to support flexible agent preferences (arbitrary weights over resources, not just dominant resource) and prioritized access (via currency burning)

while maintaining the polynomial-time solvability and starvation protection that makes DRF successful in practice. Where DRF is a heuristic that works well empirically, we provide formal theorems about optimality and fairness. Our use of specialized convex optimization (exponential cone formulations solved by Clarabel) enables us to compute exact Proportional Fairness allocations efficiently, whereas DRF uses simpler progressive filling algorithms that sacrifice optimality for speed. Where DRF treats all jobs equally, we enable priority signaling through the Priority Economy.

The practical implication is that our platform could replace or augment cluster schedulers in AI research environments, providing not just resource allocation but also fairness guarantees, collusion resistance, and audit trails. A research lab with multiple AI projects competing for GPU time could use our platform to ensure fair access while allowing urgent projects to signal priority through currency burning.

8.4 Platform Economics: The Priority Economy

The Priority Economy we introduce is a novel contribution, distinct from both traditional payment-based mechanisms and pure allocation algorithms. In traditional auctions, money is the medium of exchange: you pay to receive, you receive payment when providing. This creates closed economic loops where money circulates, wealth accumulates, and inequality compounds. In pure allocation algorithms (like DRF), there's no economic component at all—resources are allocated by formula without any notion of priority or urgency signaling.



Our Priority Economy sits between these extremes. Currency exists and agents care about it, but it doesn't circulate—it's created when resources are released and destroyed when priority is signaled. This burn-not-pay model creates fundamentally different dynamics. Agents cannot accumulate unlimited wealth because there's no way to extract value from the system. They cannot buy their way to dominance because burning currency has diminishing returns (logarithmic barrier). They must maintain economic activity (earning and spending) to maintain influence rather than sitting on accumulated wealth.

The Role of Currency

In our framework, currency is not transferred between agents or to the platform for profit. Instead, it serves as a scarce signal of urgency that modulates the arbitration objective function. This creates three interacting dynamics:

1. Earning: The "Good Citizen" Incentive

Agents earn currency primarily by voluntarily releasing resources back to the pool before their lease expires. The mechanism is simple: $\text{Income} = \text{Quantity} \times \text{TimeRemaining} \times \text{CurrentDemandMultiplier}$. An agent holding 50 compute units with 100 timesteps remaining on the lease, releasing when $\text{CurrentDemandMultiplier} = 1.5$, earns $50 \times 100 \times 1.5 = 7,500$ currency units.

This formula incentivizes exactly the behavior we want: minimize hoarding and return unused capacity immediately. Agents who hold resources longer (lower TimeRemaining) earn less per unit released. Agents who release when demand is low (lower $\text{CurrentDemandMultiplier}$) earn less. The sweet spot is releasing resources quickly when they're in demand by others—precisely the behavior that increases system liquidity and benefits everyone.

The $\text{CurrentDemandMultiplier}$ is calculated by the platform as demand/supply for each resource type, updated continuously as contention evolves. When many agents want a resource, the multiplier is high, creating strong incentives to release. When few want it, the multiplier is low, creating less pressure. This dynamic pricing naturally equilibrates supply and demand: high demand $\rightarrow$ high multiplier $\rightarrow$ agents incentivized to release $\rightarrow$ supply increases $\rightarrow$ demand decreases $\rightarrow$ multiplier falls.

2. Spending: Signaling Urgency

When an agent encounters a critical task, it "burns" its saved currency to boost its priority weight (c_i) in the Proportional Fairness optimization. The agent submits a `bid_weight` along with its resource request: "I need 50 compute units, and I'm willing to burn 100 currency to prioritize this request." The platform then solves the optimization with $c_i = \text{BaseWeight} + 100$ for that agent.

The effect is that a higher c_i forces the allocator to prioritize satisfying that agent's utility to maintain the mathematical product. But crucially, this is not pay-to-win. The logarithmic structure means that doubling your weight doesn't double your allocation—it increases it, but with diminishing returns. More importantly, no finite weight can drive another agent to zero utility because the log barrier provides infinite resistance as any agent approaches zero.

This spending mechanism enables a kind of intertemporal optimization. Agents must decide how to spend their limited currency across time: burn aggressively now to

ensure critical tasks get resources immediately, or save currency for future urgent needs? This creates interesting strategic depth while maintaining fairness constraints. The optimal strategy likely involves spending proportionally to true urgency, which means the spending mechanism elicits credible signals about which tasks are actually critical.

3. Allocation: Fairness-Mediated Priority

Money does not grant absolute authority. Unlike a "pay-to-win" auction where the highest bidder takes everything, the Weighted Proportional Fairness mechanism maintains a starvation protection guarantee through its logarithmic barrier. Because the objective function maximizes the sum of logs (equivalent to the product of utilities), the marginal gain of allocating resources to a starving agent (where $\Phi_i \rightarrow 0$) approaches infinity.

The result is that even a wealthy coalition spending massive amounts of currency cannot drive a low-resource agent's utility to zero. The mechanism allows the wealthy agent to get more or get it sooner, but never to totally starve other feasible agents. This creates a three-tier outcome space:

1. **Urgent agents with currency:** Get allocations quickly and generously
2. **Normal agents with some currency:** Get allocations at moderate speed and quantity
3. **Low-resource agents with zero currency:** Still get allocations (via BaseWeight), just slower and smaller

Crucially, all three tiers participate meaningfully. The poorest agent is not excluded, just disadvantaged. This disadvantage creates incentive to earn currency (by being a good citizen and releasing resources), which maintains system liquidity. But the disadvantage is bounded—no amount of wealth in tier 1 can push tier 3 below survival threshold.

This three-tier structure with guaranteed floor is what makes the Priority Economy fundamentally different from free markets (where the poor can be completely excluded) or command economies (where there's no signaling mechanism for urgency). It's a hybrid that combines market efficiency (urgency signaling through spending) with safety-net protection (guaranteed minimum participation). For coordinating AI agents with heterogeneous capabilities and needs, this hybrid approach provides both efficiency and safety.

8.5 Fairness in Machine Learning and Algorithmic Justice

Our work intersects with the growing literature on fairness in machine learning and algorithmic decision-making, though we address a different problem: fair resource

allocation among competing AI agents rather than fair predictions or classifications. The connection is deeper than it might initially appear—both domains grapple with how to formalize "fairness" mathematically and how to achieve it in systems where efficiency and equity potentially conflict.

Individual vs. Group Fairness

The ML fairness literature distinguishes between individual fairness (similar individuals should be treated similarly) and group fairness (protected groups should receive equal outcomes on aggregate). Dwork et al. (2012) formalized individual fairness through the notion of a task-specific similarity metric, requiring that the distance between outcomes be bounded by the distance between individuals in the input space. Our Proportional Fairness mechanism implements a form of individual fairness for agents: agents with similar preferences and priority weights receive similar allocations, with deviations justified by differences in observable preferences.

Group fairness metrics include demographic parity (groups receive equal positive classification rates), equalized odds (equal true positive and false positive rates across groups), and calibration (predictions mean the same thing across groups). Hardt et al. (2016) proved that these notions can be mutually incompatible—achieving one often precludes achieving others. Our mechanism faces an analogous tension: maximizing social welfare (efficiency) can conflict with equal allocation (equality), and our Proportional Fairness solution represents a specific resolution of this tension through logarithmic weighting.

Key difference: ML fairness typically addresses discrimination in predictive systems where a decision-maker classifies or scores individuals. Our platform addresses distribution in allocative systems where a coordinator divides scarce resources. The fairness notions differ because the underlying problems differ. Barocas et al. (2019) note that allocative harms (denying resources or opportunities) often have more immediate material impact than representational harms (stereotyping or demeaning portrayals), suggesting that getting resource allocation fairness right is particularly critical for AI safety.

Fairness-Efficiency Tradeoffs

A central question in algorithmic fairness is whether fairness must sacrifice accuracy (efficiency). Corbett-Davies and Goel (2018) argue that apparent tradeoffs often arise from inappropriate fairness metrics rather than fundamental conflicts. For resource allocation, the analogous question is whether fairness must sacrifice Pareto efficiency. Our answer is no—Proportional Fairness achieves both simultaneously (Theorem 1).

This is possible because we optimize over allocations (a continuous space with rich structure) rather than classifications (discrete decisions with limited flexibility).

However, fairness does constrain efficiency: pure utilitarian social welfare maximization (maximize $\sum_i \Phi_i$) would allocate all resources to the agent with highest marginal utility, achieving maximum total welfare but extreme inequality. By contrast, Proportional Fairness (maximize $\sum_i \log \Phi_i$) trades some total welfare for more equal distribution. The logarithmic transformation ensures that marginal utility to low-resource agents counts more heavily, preventing winner-take-all outcomes. This mirrors the fairness-accuracy tradeoff in ML: imposing fairness constraints reduces the objective value (welfare/accuracy) achievable, but prevents harmful concentration of benefits.

Kleinberg et al. (2017) showed that fairness definitions can be inherently incompatible when base rates differ across groups. For resource allocation, the analogous issue is that agents with different preferences may require different resource bundles to achieve equal utility. Our mechanism doesn't enforce equal utility (which would be inefficient) but rather equal opportunity in the form of starvation protection: all agents receive enough to maintain positive utility regardless of priority weights. This is closer to sufficiency-based fairness (ensuring minimum thresholds) than parity-based fairness (ensuring equal outcomes).

Fairness Under Uncertainty and Learning

Recent work explores fairness in dynamic settings where algorithms learn over time. Liu et al. (2018) studied how feedback loops can amplify unfairness: if an algorithm gives one group fewer opportunities, that group generates less training data, leading to worse predictions, reinforcing the disparity. Our Priority Economy faces analogous dynamics: if an agent receives poor allocations (perhaps due to low currency balance), they may fail to accomplish tasks, preventing them from earning currency through resource release, creating a poverty trap.

We address this through BaseWeight (ensuring minimum participation regardless of currency) and through the logarithmic barrier (preventing complete exclusion even temporarily). However, the formal analysis of whether these mechanisms are sufficient to prevent feedback-driven inequality remains open. Phase 2's learning component (where the platform learns to predict agent needs) introduces additional fairness risks: if the predictor learns better models for some agents than others, this could create systematic allocation disparities.

Ensign et al. (2018) proposed runaway feedback loops as a key concern in algorithmic decision systems. For our platform, the analogous risk is currency concentration spirals: wealthy agents receive better allocations, complete more tasks, earn more currency, become wealthier. The Priority Economy's bounded influence (logarithmic barrier)

provides some protection, but formal analysis of feedback stability is needed. This connects to our open problem (Section 7.3) on proving convergence of currency distributions to stable equilibria.

Algorithmic Fairness in Resource Allocation

Recent work has begun extending fairness concepts specifically to resource allocation problems. Kasy and Abebe (2021) studied optimal allocation of scarce resources (like vaccines or food aid) under fairness constraints, showing that utilitarian optimization often leads to unacceptable disparities. They proposed constrained optimization approaches similar in spirit to our weighted objective: rather than maximizing total utility subject to fairness constraints, modify the objective to internalize fairness concerns.

Our Proportional Fairness mechanism takes this approach further by making fairness the primary objective (maximize product of utilities) with efficiency emerging as a theorem (Pareto optimality) rather than an explicit constraint. This represents a philosophical shift: fairness is not a constraint on efficiency but rather a form of efficiency that accounts for distribution. Nash (1950) made this point in bargaining theory—the Nash Bargaining Solution is both fair (satisfies symmetry and scale invariance) and efficient (Pareto optimal).

Heidari et al. (2019) argued that fairness in ML should be grounded in moral and political philosophy, not just statistical notions. For resource allocation, this is even more critical: we're not just predicting outcomes but actively distributing opportunities. Our starvation protection guarantee (Theorem 3, Corollary 1.1) can be understood as implementing a Rawlsian maximin principle (prioritize the worst-off) within a Pareto-optimal framework. The logarithmic barrier creates lexicographic priority: first ensure no one starves, then maximize total welfare. This aligns with capability approaches to justice (Sen, 1999) that emphasize ensuring minimum capabilities rather than equalizing outcomes.

Connection to Our Work

Our platform-mediated coordination extends fairness considerations in several ways:

1. **Multi-dimensional resources:** ML fairness typically addresses single-dimensional decisions (hire/not hire, approve loan/deny). We handle multi-dimensional resource bundles, requiring fairness across resource types simultaneously.

2. **Strategic agents:** ML fairness assumes passive individuals receiving predictions. We deal with active agents who can adapt strategies, requiring mechanism design to prevent gaming.

3. **Temporal dynamics:** ML fairness often considers single-shot decisions. We address repeated interactions with currency dynamics and learning, requiring fairness over time.

4. **Observable-semantic gap:** ML fairness assumes the objective (accuracy, calibration) aligns with true social welfare. We explicitly model the gap between observable features and true utilities, requiring robustness to hidden preferences.

The fairness-in-ML literature provides conceptual frameworks (individual vs. group fairness, fairness-efficiency tradeoffs) that inform our mechanism design, but resource allocation for strategic agents requires new technical machinery (game theory, mechanism design, convex optimization) beyond the statistical tools of ML fairness.

8.6 Multi-Agent Reinforcement Learning

Multi-agent reinforcement learning (MARL) addresses how multiple learning agents can coordinate or compete in shared environments. While we focus on coordination through centralized arbitration rather than distributed learning, MARL provides relevant insights about strategic behavior, emergent coordination, and equilibrium concepts in multi-agent systems.

Centralized Training with Decentralized Execution (CTDE)

A prominent MARL paradigm is centralized training with decentralized execution (Kraemer and Banerjee, 2016; Oliehoek et al., 2008). During training, a centralized controller has access to all agents' states and rewards, enabling it to learn cooperative strategies. During execution, each agent acts based only on its local observations. This architecture mirrors our platform: centralized arbitration (with full observability) followed by decentralized agent execution (within allocated resources).

QMIX (Rashid et al., 2018) and QTRAN (Son et al., 2019) are CTDE algorithms that learn factored value functions: they represent the joint action-value $Q(s, a_1, \dots, a_n)$ as a mixing of individual agent values $Q_i(s_i, a_i)$. This factorization enables scalable learning while maintaining coordination. Our preference functions $\Phi_i(A)$ serve an analogous role: they factor the global allocation problem into per-agent utilities, enabling the arbitrator to reason about individual agent values while optimizing global fairness.

Key difference: MARL learns coordination policies through trial-and-error interaction, requiring many episodes to converge. Our platform computes optimal coordination through direct optimization, achieving provably Pareto-optimal allocations without learning. The tradeoff is that MARL can handle unknown or complex dynamics (learning what coordination works through experience) while our platform requires explicit preference specifications but provides stronger guarantees.

Nash Equilibria and Correlated Equilibria

MARL often analyzes learning dynamics in terms of game-theoretic equilibria. Nash equilibrium (Nash, 1951) is the most common solution concept: each agent plays a best response to others' strategies. Shoham et al. (2003) proved that many MARL algorithms converge to Nash equilibria in certain classes of games. However, Nash equilibria can be inefficient (the tragedy of the commons) and may not exist in pure strategies.

Our platform bypasses Nash equilibrium analysis by removing the game-theoretic structure entirely. Agents don't play strategies against each other; the platform directly computes optimal allocations. This is only possible because of our architectural assumption (platform visibility), but it provides stronger guarantees than Nash equilibria: Pareto optimality rather than merely mutual best-responses.

Correlated equilibria (Aumann, 1974) are a generalized solution concept where a trusted coordinator recommends actions to agents, who then best-respond to these recommendations. The correlation device can achieve outcomes unavailable to independent mixed strategies. Our platform can be understood as implementing a specific correlated equilibrium: the arbitrator "recommends" resource allocations (enforces them), and agents best-respond by using allocated resources efficiently. The difference is that we prove optimality properties (Pareto frontier) rather than merely equilibrium properties (mutual best-responses).

Emergent Communication and Cooperation

Recent MARL work studies how agents develop communication protocols to coordinate (Sukhbaatar et al., 2016; Foerster et al., 2016). Agents learn to send messages to each other, and these messages evolve to carry semantic content about intentions, observations, or proposed actions. Emergent communication enables coordination without centralized control.

Our platform provides an alternative: rather than letting agents develop communication protocols (which might be opaque, brittle, or adversarial), we provide a structured

communication channel through configuration submission. Agents "communicate" their preferences, goals, and priorities through their declared configurations, and the platform interprets this communication to compute fair allocations. This is explicit communication through a formal language rather than emergent communication through learned messages.

The advantage of explicit communication is verifiability: we can audit what agents declared and verify that allocations match these declarations. Emergent communication lacks this transparency—messages might carry semantic content that humans cannot interpret, creating safety risks as agents become more capable. For AI safety applications, the structured communication of our platform may be preferable to learned protocols.

Multi-Agent Credit Assignment

A central challenge in cooperative MARL is credit assignment: how to attribute team success or failure to individual agents' actions. COMA (Foerster et al., 2018) uses counterfactual reasoning: evaluate an agent's contribution by comparing team reward with versus without that agent's action. This requires estimating counterfactuals, which is challenging in complex environments.

Our platform's resource accounting provides explicit credit assignment: agents earn currency proportional to resources released ($\text{quantity} \times \text{remaining lease time} \times \text{demand multiplier}$). This earning mechanism attributes value precisely: an agent who releases 50 compute units when demand is high earns more than one releasing when demand is low, correctly reflecting their contribution to system liquidity. The currency ledger provides a complete audit trail of each agent's contributions and expenditures.

The Priority Economy is thus a form of market-based credit assignment: rather than trying to estimate counterfactual value through simulation, we use economic signals (earning and spending) to let value attribution emerge from agent interactions. This is simpler and more transparent than learned credit assignment but requires that value be expressible in terms of resources (which may not hold for complex cooperative tasks with emergent value).

Learning in Resource Allocation

Some MARL work specifically addresses resource allocation. Zheng et al. (2018) studied learning to allocate communication bandwidth in multi-agent systems, using RL to discover which agents should communicate when bandwidth is limited. Their approach is distributed: each agent learns a policy for requesting bandwidth. Ours is centralized: the platform computes optimal allocations.

Peysakhovich and Lerer (2018) studied prosocial learning where agents learn to cooperate through repeated interaction even without explicit communication. They found that certain RL algorithms naturally learn reciprocal strategies (cooperate with cooperators, punish defectors). Our Priority Economy creates similar dynamics: agents who "cooperate" (release resources early, creating liquidity) earn currency, enabling future priority. But this reciprocity is designed into the mechanism rather than emerged from learning.

For Phase 2 where we add learning to the platform (predicting future contentions, adapting to agent behaviors), connections to MARL become more direct. The platform could use techniques from MARL to learn models of agent dynamics: how agents react to allocations, what resources they're likely to request, when they release. This learned model would inform proactive arbitration. The key difference is that the platform is learning about agents (as a coordinator observing participants) rather than learning as an agent (as a participant in a game).

Connection to Our Work

MARL provides a contrasting paradigm to our platform approach:

MARL: Distributed learning, emergent coordination, Nash equilibria, learned communication, approximate optimality through experience

Our platform: Centralized optimization, explicit coordination, Pareto optimality, declarative communication, exact optimality through computation

The architectural difference (distributed vs. centralized) drives these distinctions. MARL is essential when you cannot assume a trusted coordinator or full observability. Our platform assumes these (viable for systems where we control the infrastructure) to achieve stronger guarantees. For AI safety, the question is: when can we design infrastructure to support platform mediation (our approach), versus when must we rely on decentralized coordination (MARL approach)?

As AI capabilities increase, the coordination challenge shifts from learning how to coordinate (solved by MARL for narrow tasks) to ensuring coordination remains safe and fair (our focus). A centralized platform can enforce properties (starvation protection, bounded inequality) that would be difficult to achieve through emergent learning alone.

8.7 Distributed Systems and Blockchain Resource Management

While our Milestone 1 implementation is centralized, the problem of coordinating resource allocation across multiple agents has been extensively studied in distributed systems and, more recently, blockchain-based resource markets. These fields provide insights about coordination without trusted central authorities, handling Byzantine faults, and using economic mechanisms to manage decentralized resources.

Distributed Resource Allocation Without Trusted Coordinators

Classical distributed systems research addresses resource allocation without centralized coordinators. Chandy-Lamport snapshot algorithm (Chandy and Lamport, 1985) enables consistent global state observation in distributed systems, analogous to our platform's omniscient state management but achieved through message passing rather than architectural visibility. Two-phase commit (Gray, 1978) and Paxos (Lamport, 1998) solve distributed consensus, relevant for a potential distributed platform where multiple arbitrator nodes must agree on allocations.

Our current architecture assumes a trusted, omniscient platform—essentially a benevolent dictator for resource coordination. Distributed systems teach us that this is a single point of failure and a potential bottleneck. A distributed version must address: How do multiple platform nodes reach consensus on allocations? How do we partition agents and resources to minimize cross-partition coordination? How do we handle network partitions where nodes disagree about resource availability?

The CAP theorem (Brewer, 2000) proves that distributed systems cannot simultaneously guarantee consistency (all nodes see the same state), availability (system always responds), and partition tolerance (system functions despite network splits). For resource allocation, this creates a fundamental tradeoff: if network partition occurs, either we maintain consistency (all nodes agree on allocations, but some agents cannot be served) or we maintain availability (all agents receive allocations, but nodes may have inconsistent resource counts, violating conservation). The design for a potential distributed version of the platform must choose where on this tradeoff to sit.

Blockchain Resource Markets

Blockchain systems provide decentralized coordination through consensus protocols and cryptographic verification. Ethereum gas markets are particularly relevant: users bid to have their transactions included in blocks, and miners select high-gas-price transactions, creating a market for limited block space (Wood, 2014). This is structurally

similar to our problem: scarce resources (block space), multiple competing agents (transaction senders), economic mechanism to allocate (gas auctions).

Key differences from our approach:

1. First-price auctions vs. Proportional Fairness: Ethereum uses first-price auctions (until recently) or EIP-1559 variable base fees, both of which are efficiency-focused mechanisms. High bidders get resources, low bidders wait or fail. There's no fairness guarantee—if you can't afford gas, your transaction doesn't happen. Our mechanism provides starvation protection (logarithmic barrier) ensuring all agents participate.

2. Pure payment vs. Priority signaling: In blockchain, you pay for resources (transaction fees go to miners). In our system, you burn currency for priority (currency is destroyed, not transferred). This changes incentives: miners are incentivized to maximize revenue, potentially colluding to inflate prices. Our platform maximizes welfare, with currency as a pure signal of urgency.

3. Permissionless vs. Platform-managed: Anyone can submit blockchain transactions. Our platform requires agent registration and configuration submission, enabling us to enforce safety invariants and detect collusion.

Bitcoin's fee market has been studied extensively as a congestion pricing mechanism (Easley et al., 2019). During high congestion, fees spike as users compete for limited block space. The market clears inefficiently: high variance in confirmation times, occasional fee overpayment (users overpay to ensure fast confirmation), and no protection for low-value but time-sensitive transactions. Our mechanism addresses these issues: predictable arbitration latency (polynomial-time optimization), no overpayment (agents don't bid against each other, they declare priorities), and protection for time-sensitive tasks (urgency signaling through currency burning).

Ethereum EIP-1559 and Mechanism Redesign

Ethereum's EIP-1559 upgrade (Buterin, 2019) replaced first-price gas auctions with a variable base fee that's burned (destroyed) plus optional priority tips for miners. This is closer to our Priority Economy: base fees are destroyed (like our currency burning) rather than transferred, and the mechanism adjusts dynamically based on demand (like our `CurrentDemandMultiplier` in earning formula).

Similarities:

- Currency burning as a coordination signal (not wealth transfer)
- Dynamic adjustment to congestion levels

- Separation of base cost (base fee) from priority (tips)

Differences:

- EIP-1559 optimizes for stable fees and block utilization, not fairness
- No starvation protection (if you can't afford base fee, you don't transact)
- Purely reactive (fees adjust after congestion) rather than predictive

The EIP-1559 design space exploration (Roughgarden, 2020) analyzed tradeoffs between incentive compatibility, miner revenue, user experience, and system stability. For our Priority Economy, analogous analysis is needed: tradeoffs between fairness guarantees, economic efficiency, strategic manipulability, and currency equilibrium stability. Section 7.3 identifies this as open research.

Decentralized Resource Allocation Protocols

Beyond blockchains, distributed systems research has explored decentralized resource allocation without centralized markets. Resource allocation in peer-to-peer systems (BitTorrent's tit-for-tat, Gnutella's freeloader problem) demonstrates challenges of fairness without central authority: tragedy of the commons, free-riding incentives, difficulty of punishing bad actors.

Toka et al. (2009) proposed game-theoretic mechanisms for fair resource allocation in peer-to-peer systems, using reputation and reciprocity. Their approach is purely decentralized: agents directly negotiate with peers, building reputation through cooperation. This contrasts with our platform approach but offers resilience: no single point of failure, no central authority to trust or attack. The tradeoff is weaker fairness guarantees (reputation can be gamed) and higher complexity (agents must maintain state about peers).

For AI safety, the question is architectural: Should we aim for decentralized coordination (resilient but weaker guarantees) or centralized platforms (strong guarantees but central points of failure)? Our current work assumes centralized is feasible and preferable, but as systems scale or span trust boundaries, decentralized approaches may become necessary.

Byzantine Fault Tolerance

Byzantine agents (malicious, arbitrary behavior) are extensively studied in distributed systems. PBFT (Castro and Liskov, 1999) achieves consensus with up to $f < n/3$ Byzantine nodes through cryptographic signatures and voting. For our platform,

Byzantine agents could: report false utilities, attempt to violate resource conservation through double-spending, or collude to disrupt arbitration.

Milestone 1 assumes non-Byzantine agents (agents may strategize but not actively sabotage). Future phases must address:

1. **Cryptographic resource accounting:** Sign all resource transfers, making double-spending detectable
2. **Byzantine consensus for arbitration:** If multiple arbitrator nodes exist, require $2f+1$ agreement on allocations
3. **Agent quarantine:** Detect Byzantine behavior through anomaly detection, isolate malicious agents
4. **Stake-based admission:** Require agents to stake currency upon registration, slashing stake if Byzantine behavior detected

The Ripple consensus mechanism (Schwartz et al., 2014) uses a trusted subgraph approach: each node specifies which other nodes it trusts, and consensus requires agreement among trusted nodes. For our platform, this could mean: agents specify which other agents they trust (won't collude with), and the arbitrator ensures allocations are fair to trust networks. This is speculative (Phase 4+) but suggests how trust structures could augment mechanism design.

Connection to Our Work

Distributed systems and blockchain research shows that decentralized coordination is possible but requires careful incentive design and accepts weaker guarantees than centralized approaches. Our platform takes the centralized route because:

1. We control the infrastructure (agents execute within our sandbox)
2. We prioritize strong guarantees (Pareto optimality, starvation protection) over decentralization
3. We target safety-critical applications where trusted arbitration is acceptable

However, as scale increases ($n > 1000$ agents in contention) or trust boundaries emerge (coordinating across organizations), lessons from distributed systems become critical. Any potential distributed architecture will likely borrow from blockchain consensus (Raft for consistency) and peer-to-peer resource allocation (partitioning strategies, cross-partition protocols).

The key insight from blockchain resource markets is that pure economic mechanisms (auctions, fees) without fairness constraints lead to winner-take-all outcomes.

EIP-1559's burning mechanism shows that destroying currency rather than transferring it can improve incentives, but even EIP-1559 lacks our starvation protection. For AI coordination as capabilities increase, fairness guarantees become more critical, suggesting that pure market mechanisms are insufficient—we need mechanisms like Proportional Fairness that internalize fairness into the objective.

8.8 Cooperative Game Theory and Social Choice

While we've mentioned Nash bargaining, the broader literature on cooperative game theory and social choice theory provides foundational concepts for thinking about fair allocation and collective decision-making that inform our mechanism design.

The Core and Stable Allocations

Cooperative game theory studies coalitions: groups of agents who can jointly deviate from proposed outcomes. The core is the set of allocations where no coalition can improve by deviating (Gillies, 1959). An allocation is in the core if it's stable—no subset of agents would prefer to "leave" and divide resources among themselves.

Our Proportional Fairness mechanism does not explicitly guarantee core membership. In fact, for resource allocation problems, the core may be empty (no stable allocation exists) or very small. However, the logarithmic barrier provides a form of stability: coalitions attempting to monopolize resources face diminishing returns and resistance from the mechanism's defense of victims (Theorem 3). This is weaker than core stability but stronger than nothing—it's a form of "approximate stability" where deviations are unprofitable due to mechanism structure rather than pure coalition incentives.

Shapley value (Shapley, 1953) is another solution concept: it assigns each agent a value equal to their average marginal contribution across all possible coalition formations. For resource allocation, this would mean: agent i 's fair share equals the average of how much they improve allocations across all possible subsets they could join. Computing Shapley values is exponentially hard (requires enumerating 2^n coalitions), making it impractical for our scale. However, Shapley value provides a normative benchmark: does our Proportional Fairness allocation approximate Shapley values? This could be tested empirically in Milestone 2.

Envy-Freeness and Proportionality

Fair division literature defines multiple fairness notions (Brams and Taylor, 1996):

Envy-freeness: No agent prefers another agent's allocation to their own. For heterogeneous preferences, this is very strong—it requires that agents with different preferences all simultaneously feel they received the best bundle for themselves.

Proportionality: Each of n agents receives at least $1/n$ of their value. Weaker than envy-freeness but still demanding.

Proportional fairness (our mechanism): Maximizes product of utilities, equivalent to Nash bargaining solution. Related to but distinct from proportionality above.

Our mechanism doesn't guarantee envy-freeness. Agent A_1 with $w_{\text{compute}} = 0.8$ might envy Agent A_2 's allocation if A_2 received 60 compute units while A_1 received only 50 (even if A_1 's total utility is higher due to other resources). This is acceptable—envy-freeness is often too strong for heterogeneous preferences, requiring Pareto-inefficient allocations to eliminate all envy.

Proportionality is closer to what we achieve. The logarithmic barrier ensures no agent receives utility too far below their "fair share," though defining "fair share" is subtle when preferences differ. One interpretation: each agent should receive enough utility that their log term contributes meaningfully to the objective. Since $\sum_i c_i \log(\Phi_i)$ is maximized, agents with very low Φ_i would have very negative log terms, dragging down the objective. The mechanism prevents this by allocating enough to keep all Φ_i reasonably large.

Impossibility Results and Tradeoffs

Arrow's impossibility theorem (Arrow, 1950) proves that no voting system can simultaneously satisfy: universality (works for all preference profiles), non-dictatorship (no single voter determines outcome), Pareto efficiency, and independence of irrelevant alternatives (adding candidates doesn't change relative rankings). This is devastating for social choice: any voting mechanism must violate at least one desirable property.

For resource allocation, analogous impossibility results exist. Moulin (1990) proved that no strategyproof mechanism can be both Pareto-efficient and envy-free for economies with multiple goods. Our mechanism "escapes" this impossibility by not requiring strategyproof reporting—we directly observe preferences through platform visibility, avoiding the strategic reporting problem that creates impossibility.

However, we face a different impossibility: we cannot simultaneously achieve exact Pareto optimality, complete collusion-proofness, and computational tractability for

arbitrary nonlinear preferences. Theorem 2 proves tractability for linear preferences. Theorem 3 argues (heuristically) for collusion resistance. Theorem 1 proves Pareto optimality. But extending to arbitrary nonlinear utilities (Phase 2) may force us to sacrifice one of these properties—likely requiring approximation for computational tractability.

Voting Rules and Aggregation

Social choice studies how to aggregate individual preferences into collective decisions. Voting rules (plurality, Borda count, Condorcet methods) can be understood as ways to allocate a single resource (who wins the election). Our mechanism allocates multiple resources, but the principle is similar: how do we aggregate agents' individual preferences (w_{ij}) into a collective allocation decision?

Quadratic voting (Posner and Weyl, 2018) is a recent proposal where voters can buy votes at quadratic cost: buying n votes costs n^2 currency. This creates diminishing returns to vote buying, preventing wealthy voters from completely dominating. Our Priority Economy has a similar flavor: burning currency increases priority weight c_i , but the logarithmic objective means marginal gains diminish. Burning $2x$ currency doesn't give $2x$ allocation—the relationship is sublinear.

However, quadratic voting achieves incentive compatibility (truthful reporting is optimal) through its cost structure. Our mechanism doesn't require incentive compatibility (we observe preferences directly) but faces a different challenge: coalitions coordinating currency burning. Quadratic voting's resistance to collusion comes from its cost structure; ours comes from the logarithmic barrier. Both mechanisms make concentrated power expensive through mathematical structure rather than through external enforcement.

Justice and Fairness Theories

Philosophical theories of distributive justice inform what "fairness" should mean. Rawls' difference principle (Rawls, 1971) says inequalities are justified only if they improve the worst-off agent's position. This is stronger than Pareto optimality (which allows inequalities even if they don't help the worst-off) but weaker than egalitarianism (which eliminates all inequality regardless of efficiency).

Our logarithmic barrier implements a softer version of Rawls: we don't strictly maximize the minimum utility (maximin), but we heavily weight improvements to low-utility agents through the log transformation. An agent at $\Phi_i = 1$ has marginal weight $c_i/1 = c_i$. An agent at $\Phi_i = 100$ has marginal weight $c_i/100$. The mechanism implicitly prioritizes the former 100x more heavily, creating Rawlsian-flavored outcomes without pure maximin.

Sen's capability approach (Sen, 1999) argues fairness should ensure capabilities (what people can do) rather than resources (what they have) or welfare (how satisfied they are). For our platform, this suggests that minimum thresholds on resource allocations (ensuring agents can accomplish their goals) are more important than equalizing utilities. Our goal satisfaction framework (Section 2.2, goals Γ) captures this: agents specify what they need to function, and the platform ensures these minimums are respected when feasible.

Connection to Our Work

Cooperative game theory and social choice provide normative frameworks for evaluating our mechanism:

What we achieve: Pareto optimality (Theorem 1), approximate stability (Theorem 3's collusion resistance), bounded inequality (logarithmic barrier), computability (Theorem 2).

What we don't achieve: Core stability (coalitions might have incentive to deviate), envy-freeness (agents might prefer others' allocations), strategyproofness (not needed due to platform visibility, but worth analyzing for hidden utilities).

Open questions: Does our allocation approximate Shapley values? Under what conditions is it in the core? How does it compare to other fairness notions (maximin, leximin, proportionality) empirically?

These questions don't undermine the mechanism—Proportional Fairness is well-founded in axiomatic bargaining theory—but they suggest fruitful directions for deeper analysis. For AI safety applications, understanding which fairness properties we provide (and which we don't) clarifies what behaviors to expect as agent capabilities increase.

8.9 Comparative Analysis of Coordination Mechanisms (Expanded)

To contextualize our contribution, we provide systematic comparison of our Proportional Fairness mechanism against established alternatives across key properties relevant to AI safety and multi-agent coordination. This expanded analysis includes mechanisms

from distributed systems, game theory, and resource allocation beyond our original comparison.

8.9.1 Mechanism Properties Comparison

Table 1. Comparative Analysis of Coordination Mechanisms against AI Safety Requirements

Mechanism	Pareto Optimal	Collusion Resistant	Polynomial Time	Starvation Protection	Incentive Compatible	Scalability	Implementation Complexity
Proportional Fairness (Ours)	✓ (proven)	✓ (heuristic)	✓ (proven)	✓ (logarithmic barrier)	N/A (direct observation)	$n \leq 50$	Medium
VCG/Groves	✓	✗ (vulnerable)	✗ (NP-hard)	✗	✓ (dominant strategy)	$n \leq 20$	High
Dominant Resource Fairness (DRF)	✗ (heuristic)	~ (limited)	✓ (progressive filling)	✓ (proportional)	N/A (centralized)	$n \leq 1000$	Low
Equal Allocation	✗	✓	✓ (trivial)	✓	N/A	$n \leq \infty$	Very Low
Proportional to Requests	✗	✗ (request inflation)	✓ (linear)	~ (depends on requests)	✗ (underreporting)	$n \leq \infty$	Low
First-Come-First-Served	✗	~ (timing games)	✓ (queue processing)	✗ (latecomers starve)	✗ (early request incentive)	$n \leq \infty$	Low

Mechanism	Pareto Optimal	Collusion Resistant	Polynomial Time	Starvation Protection	Incentive Compatible	Scalability	Implementation Complexity
Nash Bargaining	✓ (same as PF)	~ (threat point dependent)	? (formulation dependent)	~ (disagreement utility)	~ (protocol dependent)	$n \leq 10$	High
Blockchain Gas Auctions	✗ (efficiency only)	~ (fee manipulation)	✓ (memory processing)	✗	✓ (first-price IC)	$n \leq \infty$	Medium
EIP-1559 (Ethereum)	~ (block utilization)	~ (fee burning helps)	✓	✗	~ (strategy-robust)	$n \leq \infty$	Medium
Quadratic Voting	~ (single-item only)	✓ (quadratic cost)	✓	~ (depends on budget)	✓ (truthful bidding)	$n \leq 1000$	Medium
MARL (Nash Equilibrium)	✗ (equilibrium $\neq$ optimal)	✗ (collusion equilibria)	✗ (learning time)	✗	~ (learned strategies)	$n \leq 100$	High
Shapley Value	✓	✓ (axiomatically fair)	✗ (exponential)	✓ (contribution-based)	✓	$n \leq 20$	Very High
Core Allocation	✓ (by definition)	✓ (stable)	✗ (exponential)	~ (depends on coalition structure)	✓	$n \leq 15$	Very High

Mechanism	Pareto Optimal	Collusion Resistant	Polynomial Time	Starvation Protection	Incentive Compatible	Scalability	Implementation Complexity
Maximin (Rawlsian)	✗ (sacrifices efficiency)	✓	✓	✓ (maximum protection)	N/A	$n \leq \infty$	Low
Utilitarian Allocation	✓	✗ (winner-take-all)	✓	✗	N/A	$n \leq \infty$	Low

Legend:

- ✓ = Property holds with formal proof or strong evidence
- ~ = Property partially holds or depends on parameters
- ✗ = Property does not hold or explicitly vulnerable
- N/A = Property not applicable to mechanism design

8.9.2 Detailed Mechanism Analysis

VCG (Vickrey-Clarke-Groves) Mechanisms

Mechanism: Allocate resources to maximize social welfare $\sum_i v_i(\text{allocation})$, charge each agent their externality (harm to others from their participation).

Advantages:

- **Gold standard for incentive compatibility:** Truthful reporting is a dominant strategy
- Works with agents as black boxes (doesn't require platform visibility)
- Extensive theoretical foundation (60+ years)
- Guarantees efficiency (maximizes social welfare)

Disadvantages:

- **Computational vulnerability:** Optimal allocation requires solving NP-hard combinatorial problems

- **Collusion vulnerability:** Coalitions can profitably manipulate through shill bidding, demand reduction
- **No fairness guarantees:** Maximizes total welfare but can produce extreme inequality (one agent gets everything)
- **Payment complexity:** Requires monetary transfers, budget constraints, revenue distribution
- **Participation issues:** Individual rationality not always satisfied (agents might lose money)

When suitable: Small-scale systems ($n \leq 20$) with truly adversarial agents who cannot be sandboxed, where incentive compatibility is paramount and computational cost is acceptable.

Why we chose differently: For platform-mediated AI coordination where visibility is architectural, trading VCG's incentive properties for computational security and fairness guarantees is worthwhile. Attack surface from NP-hard computation becomes increasingly dangerous as agents grow sophisticated.

Dominant Resource Fairness (DRF)

Mechanism: Allocate each agent's "dominant share" (their maximum proportional share across resource types) equally. Agent i 's dominant resource is $\arg\max_j (\text{alloc}_{ij} / \text{total}_j)$.

Advantages:

- **Simple and fast:** Progressive filling algorithm, $O(nm \log n)$ complexity
- **Proven at scale:** Production deployment in Hadoop YARN, Apache Mesos (1000s of machines)
- **Strategy-proof** (in basic formulation): Agents cannot gain by misreporting
- **Strong fairness properties:** Share guarantee, envy-freeness, Pareto efficiency within dominant resource paradigm

Disadvantages:

- **Not globally Pareto optimal:** Only ensures Pareto efficiency within dominant resource framework, misses cross-resource trades
- **Limited preference expressiveness:** Only considers dominant resource, ignores agent weights over other resources
- **No priority mechanism:** All agents treated equally, cannot signal urgency
- **Resource-type specific:** Designed for computational resources, unclear how to extend to physical/informational resources

When suitable: High-scale cluster scheduling ($n > 100$) where simplicity and speed are paramount, agents have clear dominant resource needs, approximate optimality is acceptable.

Why we chose differently: DRF's heuristic nature leaves Pareto improvements unexploited. Our research contribution is showing exact optimality is achievable at similar computational cost using convex optimization, providing stronger guarantees while maintaining practical performance.

Equal Allocation

Mechanism: Divide resources equally among all agents: $\text{alloc}_{i,j} = \text{total}_j / n$ for all i, j .

Advantages:

- **Perfect equality:** No inequality by construction, complete fairness
- **Trivially simple:** No computation beyond division
- **No strategic manipulation:** Nothing to gain from misreporting

Disadvantages:

- **Ignores preferences:** Agent who needs compute gets same amount as agent who doesn't care
- **Extremely inefficient:** Massive Pareto improvements left unexploited (agent who values compute could trade with agent who values storage)
- **Doesn't handle heterogeneity:** Agents with different scales or needs treated identically

When suitable: Small, homogeneous systems where equality is paramount and efficiency is secondary, or as a fallback when preference information is unavailable.

Why we chose differently: Equal allocation is a useful baseline but sacrificing all efficiency for perfect equality is untenable for AI coordination where agents have genuinely different goals and resource needs.

Proportional Allocation (to Requests)

Mechanism: Allocate proportionally to requests: if agent i requests $r_{i,j}$ of resource j and total demand $\sum_i r_{i,j}$ exceeds supply, allocate $\text{alloc}_{i,j} = r_{i,j} \cdot (\text{total}_j / \sum_i r_{i,j})$.

Advantages:

- **Simple linear formula:** Easy to compute and explain

- **Intuitively "fair"**: Everyone gets proportional share
- **Handles partial allocation**: Naturally reduces requests when demand exceeds supply

Disadvantages:

- **Not Pareto optimal**: Ignores preferences—agent desperately needing resource gets same proportion as agent mildly wanting it
- **Strategic vulnerability**: Agents incentivized to inflate requests to get larger proportional shares
- **No starvation protection**: Agent requesting 1000x more than others gets 1000x more, monopolizing resources
- **Request explosion**: In equilibrium, all agents request maximum to ensure proportional share

When suitable: Systems where preferences are unknown and requests are somewhat honest, or as a simple heuristic when optimization is infeasible.

Why we chose differently: Proportional allocation is a common practical heuristic but lacks both efficiency guarantees (Pareto optimality) and fairness guarantees (starvation protection) needed for AI safety applications.

First-Come-First-Served (FCFS)

Mechanism: Process resource requests in temporal order until resources are exhausted.

Advantages:

- **Simplest scheduling**: Just maintain a queue
- **Natural temporal priority**: Earlier requests served first
- **Easy to understand**: Transparent, predictable

Disadvantages:

- **Ignores preferences and urgency**: All requests treated identically
- **Vulnerable to timing games**: Agents submit requests speculatively or repeatedly to ensure early position
- **Starvation for late arrivals**: Once resources consumed, latecomers receive nothing
- **No fairness properties**: Completely arbitrary outcomes based on timing accidents

When suitable: Single-resource systems with genuinely time-ordered requests where temporal priority is meaningful, or embedded systems with hard real-time constraints.

Why we chose differently: FCFS is inappropriate for multi-agent coordination where fairness matters. For AI safety, we cannot accept mechanisms where timing accidents determine resource access.

Nash Bargaining Solution

Mechanism: Maximize product of utilities (same as Proportional Fairness) subject to individual rationality and Pareto optimality, with disagreement point defining outside options.

Relationship to our work: Our Proportional Fairness IS Nash Bargaining applied to resource allocation.

What we add to Nash:

1. **Computational implementation:** Convex optimization formulation with polynomial-time guarantees
2. **Priority weighting:** Currency burning extends Nash to weighted bargaining
3. **Multi-agent generalization:** Nash originally for 2-player bargaining, we handle n-agent
4. **Platform architecture:** Practical system design for executing Nash solution
5. **Safety analysis:** Collusion resistance, starvation protection, semantic divergence bounds

Distinction: Nash proved uniqueness and axioms (1950), Kelly applied to networks (1997), we apply to AI coordination with computational security and safety guarantees (2024).

Blockchain Gas Auctions

Mechanism: Users bid gas price (per computational step), miners select high-gas transactions to maximize revenue.

Advantages:

- **Permissionless:** Anyone can submit transactions
- **Decentralized:** No central coordinator or single point of failure
- **Responsive:** Prices adjust quickly to demand
- **Incentive-aligned:** Miners profit from efficient resource use

Disadvantages:

- **No fairness guarantees:** High bidders get resources, low bidders wait indefinitely or fail
- **Price volatility:** Gas prices can spike 100x during congestion
- **Poor user experience:** Confirmation time unpredictable, frequent overpayment
- **Miner extractable value (MEV):** Miners can manipulate transaction ordering for profit
- **Starvation:** Low-value but time-sensitive transactions starve during congestion

When suitable: Permissionless systems where decentralization is paramount, users have heterogeneous valuations, and fairness is not a concern.

Why we chose differently: For coordinating AI agents within controlled infrastructure, centralized arbitration enables fairness guarantees unavailable in decentralized markets. Starvation protection is critical for AI safety.

EIP-1559 (Ethereum Fee Market)

Mechanism: Variable base fee (burned) that adjusts to target 50% block utilization, plus optional priority fees (paid to miners).

Advantages:

- **Improved UX:** More predictable fees than pure auctions
- **Burn mechanism:** Base fees destroyed, reducing ETH supply, aligning incentives
- **Congestion control:** Base fee adjusts to manage demand
- **Separation:** Base cost separate from priority, clarifying what fees pay for

Disadvantages:

- **Still no fairness:** If you can't afford base fee, you don't transact
- **Reactive:** Fees adjust after congestion occurs, not predictive
- **Priority fees create inequality:** Wealthy users can still skip queue
- **No starvation protection:** Low-value transactions excluded during high demand

Relationship to our work: Priority Economy similarly uses burning (currency destroyed, not transferred) to signal urgency. But we add fairness (logarithmic barrier) while EIP-1559 optimizes for stable fees and block utilization.

Quadratic Voting

Mechanism: Voters can buy votes at quadratic cost: n votes cost n^2 currency. Applied to resource allocation: agent can buy priority at quadratic cost.

Advantages:

- **Collusion resistance:** Quadratic cost makes vote buying expensive for coalitions
- **Incentive compatible:** Truthful voting (voting proportional to intensity) is optimal
- **Efficient:** Achieves approximate social welfare maximization
- **Formal analysis:** Weyl and Posner prove optimality properties under assumptions

Disadvantages:

- **Single-item focus:** Designed for voting on single alternatives, unclear how to extend to multi-resource allocation
- **Budget dependency:** Requires agents have comparable budgets, else wealthy dominate
- **Complex explanation:** Quadratic cost structure is non-intuitive
- **No Pareto optimality:** Maximizes welfare but doesn't ensure no waste

Relationship to our work: Both mechanisms make concentrated power expensive through mathematical structure (quadratic cost vs. logarithmic barrier). Both use currency non-transferably (votes are bought/spent, not traded). Difference: quadratic voting for single decisions, our mechanism for multi-resource allocations.

MARL (Multi-Agent Reinforcement Learning) with Nash Equilibrium

Mechanism: Agents learn policies through trial-and-error, converging to Nash equilibrium where each agent best-responds to others.

Advantages:

- **Decentralized:** No coordinator needed
- **Adaptive:** Learns from experience, handles unknown dynamics
- **Scalable to large state spaces:** Deep RL handles high-dimensional problems
- **Emergent coordination:** Can discover complex cooperative strategies

Disadvantages:

- **Not Pareto optimal:** Nash equilibria can be arbitrarily inefficient (e.g., prisoner's dilemma)
- **Collusion equilibria:** Coalitions can learn to coordinate exploitation
- **Slow convergence:** Requires many episodes to learn
- **No guarantees:** May not converge, may converge to bad equilibrium
- **Opaque:** Learned policies difficult to interpret or verify

When suitable: Decentralized systems with unknown dynamics, where learning is necessary and approximate coordination acceptable.

Why we chose differently: For safety-critical AI coordination where we control infrastructure, direct optimization provides stronger guarantees than learned equilibria. Verifiability requires explicit mechanisms, not emergent strategies.

Shapley Value

Mechanism: Allocate each agent value equal to their average marginal contribution across all possible coalition orderings.

Advantages:

- **Axiomatically fair:** Satisfies efficiency, symmetry, dummy player, additivity
- **Coalitionally stable:** In the core for convex games
- **Contribution-based:** Rewards agents proportional to value they add
- **Unique:** Only solution satisfying Shapley's axioms

Disadvantages:

- **Computationally intractable:** Requires evaluating 2^n coalitions, exponential in n
- **Requires value function:** Need to evaluate worth of every coalition, often unknown
- **Not obviously optimal:** Satisfies fairness axioms but not Pareto optimality axioms

When suitable: Small cooperative games ($n \leq 10$) where coalition values are known, fairness according to Shapley axioms is desired, computational cost is acceptable.

Why we chose differently: Exponential complexity makes Shapley value infeasible for $n > 20$. Our mechanism achieves different fairness notion (Proportional Fairness, not Shapley) in polynomial time.

Open question: Does our allocation empirically approximate Shapley values? Could test in Milestone 2 for small n .

Core Allocation

Mechanism: Find allocation where no coalition can improve by deviating—no subset of agents would prefer to "leave" with their resources.

Advantages:

- **Coalitionally stable:** By definition, no profitable deviations
- **Pareto optimal:** Core allocations are always efficient
- **Strong fairness:** Satisfies many intuitive fairness notions

Disadvantages:

- **May not exist:** Many games have empty core (no stable allocation)
- **Computationally intractable:** Finding core allocations requires checking exponentially many coalitions
- **Underdetermined:** Core typically contains many allocations, need tiebreaker

When suitable: Cooperative settings where coalition stability is paramount, games known to have non-empty core, small n where computation feasible.

Why we chose differently: Core membership is too strong a requirement (may not exist) and too expensive to compute. Our logarithmic barrier provides approximate stability (coalitions face diminishing returns) without requiring exact core membership.

Maximin (Rawlsian) Allocation

Mechanism: Maximize minimum utility: $\max_{\text{allocation}} \min_i \Phi_i(\text{allocation})$. Prioritizes worst-off agent.

Advantages:

- **Maximum starvation protection:** Strongest possible guarantee for worst-off
- **Philosophically grounded:** Implements Rawls' difference principle
- **Robust:** Worst-case optimization naturally handles uncertainty

Disadvantages:

- **Sacrifices efficiency:** May leave large Pareto improvements unexploited to marginally help worst-off
- **Severe inequality:** To maximize minimum, may create huge gaps between worst-off and best-off
- **Ignores preference intensity:** Treats all increases to minimum utility equally regardless of cost

When suitable: Systems prioritizing absolute protection of vulnerable agents above efficiency, disaster relief or crisis response scenarios.

Why we chose differently: Pure maximin is too rigid—we want starvation protection (logarithmic barrier) but not at the expense of all efficiency (Pareto optimality). Proportional Fairness balances protection and efficiency.

Utilitarian (Social Welfare Maximization)

Mechanism: Maximize sum of utilities: $\max_{\text{allocation}} \sum_i \Phi_i(\text{allocation})$. Pure efficiency focus.

Advantages:

- **Maximum efficiency:** Achieves highest total welfare
- **Simple objective:** Easy to understand and compute
- **Pareto optimal:** By construction

Disadvantages:

- **No fairness:** Can create extreme inequality (allocate everything to agent with highest marginal utility)
- **No starvation protection:** Agents with low values may receive nothing
- **Ignores distribution:** Total welfare could be 100 with one agent getting 100 and others getting 0

When suitable: Settings where efficiency is sole concern, agents are homogeneous, or redistribution can be handled separately.

Why we chose differently: For AI safety, distribution matters. Mechanisms that allow complete exclusion create dangerous power concentration. We need efficiency (Pareto optimality) AND fairness (starvation protection), which Proportional Fairness provides.

8.9.3 Suitability for AI Safety Applications

Different coordination mechanisms are appropriate for different contexts. For AI safety applications involving coordination of increasingly capable AI agents where catastrophic failure modes must be prevented, certain properties become critical rather than merely desirable.

Critical Requirements for Safe AI Coordination

As AI capabilities approach and potentially exceed human-level performance across domains, coordination mechanisms must satisfy stringent requirements:

1. Starvation protection: No agent can be completely excluded from resources regardless of others' power or coordination. This prevents winner-take-all dynamics where dominant agents or coalitions monopolize resources, creating dangerous power concentration and eliminating system diversity that provides resilience.

2. Collusion resistance: Coordinated manipulation by coalitions must be structurally unprofitable or detectable. As AI agents become more capable of sophisticated coordination, mechanisms relying on competitive friction to prevent collusion will fail. Mathematical structure that makes collusion self-defeating provides more robust protection.

3. Computational security: No exponential worst-case complexity that enables denial-of-service attacks. Sophisticated agents could craft inputs that trigger pathological algorithm behavior, effectively freezing coordination. Polynomial-time guarantees ensure adversaries cannot weaponize computational complexity.

4. Verifiability: Allocations must be auditable and provable against mechanism specification. As AI systems make consequential decisions, we need to verify that coordination mechanisms function as designed rather than relying on opacity and hope. Deterministic, explicit algorithms enable verification; learned policies and opaque heuristics do not.

Secondary requirements include efficiency (Pareto optimality or close approximation), bounded inequality (preventing extreme wealth/power concentration), and scalability (handling growth in agent populations).

Mechanism Suitability Assessment

Based on the detailed analysis in Section 8.9.2 and the comparison table in Section 8.9.1, we can assess which mechanisms meet these critical requirements:

Mechanisms satisfying all four critical requirements:

- **Proportional Fairness (our mechanism):** Logarithmic barrier provides starvation protection (Theorem 3), structural collusion resistance through hypersensitivity to victim utility approaching zero, polynomial-time computation via convex optimization (Theorem 2), and full verifiability through deterministic optimization. Additionally achieves Pareto optimality (Theorem 1) and bounds inequality through logarithmic weighting. Scalability limitation ($n \leq 50$ size contentions for Milestone 1) is acceptable for initial deployment.

- **Shapley Value:** Axiomatic fairness provides starvation protection (all contributors receive positive value), collusion resistance through contribution-based allocation, full verifiability through explicit formula. However, exponential computational complexity (2^n coalitions) limits practical deployment to $n \leq 10$, making it unsuitable for most AI coordination scenarios despite excellent theoretical properties.

Mechanisms satisfying three critical requirements:

- **DRF (Dominant Resource Fairness):** Provides starvation protection through proportional share guarantees, polynomial-time computation, and full verifiability. Moderate collusion resistance (proportional sharing makes manipulation less profitable but not impossible). Sacrifices exact Pareto optimality (achieves approximate efficiency within dominant resource framework). Excellent scalability ($n \geq 1000$) makes it attractive when scale is paramount and approximate optimality is acceptable.

- **Equal Allocation:** Trivially satisfies computational security, verifiability, starvation protection, and collusion resistance (nothing to manipulate). However, catastrophically inefficient (ignores all preference information), making it suitable only as a fallback when no preference information is available.

Mechanisms satisfying two or fewer critical requirements:

- **VCG:** Achieves only computational security (fails in worst case) and verifiability. Lacks starvation protection (can allocate everything to highest-value agent) and is highly vulnerable to collusion (bidding rings, shill bidding). Despite being the "gold standard" for incentive compatibility in mechanism design, VCG is unsuitable for AI safety coordination where fairness and collusion resistance are critical.

- **Market mechanisms (Gas Auctions, EIP-1559):** Achieve computational security and verifiability but lack starvation protection (agents who cannot afford fees are excluded) and collusion resistance (fee manipulation, miner extractable value). Suitable for permissionless systems prioritizing decentralization over fairness, unsuitable for coordinating powerful AI agents where exclusion creates dangerous power concentration.

- **MARL (Nash Equilibrium):** Fails all four critical requirements. Nash equilibria provide no efficiency guarantee (can be arbitrarily suboptimal), no fairness guarantee (can completely exclude agents), no collusion resistance (collusive equilibria exist), and poor verifiability (learned policies are opaque). Suitable for decentralized systems learning coordination in unknown environments, unsuitable for safety-critical applications requiring provable properties.

- **Maximin**: Provides maximum starvation protection (optimizes for worst-off agent) but at severe efficiency cost. Computational security and verifiability are satisfied. However, extreme inequality between worst-off and best-off agents can emerge (everything goes to slightly improving minimum). Suitable for disaster response or crisis scenarios where protecting most vulnerable is paramount, less suitable for ongoing coordination where efficiency matters.

- **Utilitarian (Pure Social Welfare)**: Achieves Pareto optimality by construction but fails starvation protection (can allocate everything to highest marginal utility agent), collusion resistance (coalitions can coordinate to exclude others), and creates extreme inequality. Suitable only when fairness is handled through separate redistribution mechanisms.

Architectural Considerations

The suitability analysis must also consider architectural assumptions each mechanism requires:

Centralized with full observability (our approach): Proportional Fairness, DRF, Equal Allocation, Maximin, Utilitarian all assume a trusted coordinator with complete visibility into agent states and preferences. This assumption is viable when we control the infrastructure (agents execute within platform sandbox) but may not hold when coordinating across organizational or trust boundaries.

Decentralized with limited observability: MARL, blockchain mechanisms (Gas Auctions, EIP-1559), and peer-to-peer protocols operate without central coordination. This provides resilience (no single point of failure) but achieves weaker guarantees. For AI safety, the question is: when can we design infrastructure to support centralized arbitration (enabling stronger guarantees), versus when must we rely on decentralized coordination (accepting weaker guarantees for resilience)?

Strategic reporting vs. direct observation: VCG and Quadratic Voting solve the preference elicitation problem through clever incentive design (making truthful reporting optimal). Our platform bypasses this problem through architectural visibility—we observe preferences directly from configurations rather than asking agents to report them. This shift eliminates the strategic reporting problem but introduces a different challenge: the observable-semantic utility gap (Theorem 4), where agents might hide utility components the platform cannot observe.

Scaling and Phase Transitions

Mechanism suitability changes as coordination scale increases:

Small scale ($n \leq 20$): All mechanisms except Shapley Value are computationally feasible. Choice depends on which properties matter most. For AI safety research labs coordinating GPU access among projects, Proportional Fairness provides optimal balance of fairness, efficiency, and security.

Medium scale ($n \leq 100$): VCG becomes computationally risky (exponential worst-case), Shapley Value infeasible. Proportional Fairness requires distributed architecture or approximation. DRF excels in this range, making it attractive for coordinating AI agent fleets within organizations.

Large scale ($n \geq 1000$): Only mechanisms with near-linear complexity remain practical (Equal Allocation, Proportional to Requests, FCFS, DRF, market mechanisms). For global AI coordination across many organizations, tradeoffs become unavoidable: either sacrifice fairness guarantees (market mechanisms) or sacrifice efficiency (DRF's approximation). This suggests large-scale coordination may require hierarchical approaches: Proportional Fairness within organizations (providing strong fairness guarantees), market-like mechanisms between organizations (accepting weaker guarantees for scalability).

Capability Level Considerations

The importance of different properties shifts as AI agent capabilities increase:

Narrow AI (task-specific systems): Strategic behavior is limited, collusion requires explicit coordination by creators. Simple mechanisms (FCFS, proportional allocation) may suffice. However, establishing norms of fair coordination now creates precedents for future more capable systems.

Near-term advanced general-purpose AI (multi-domain capable agents): Strategic behavior becomes sophisticated, collusion easier to coordinate. Starvation protection becomes critical as capability disparities increase. Mechanisms like Proportional Fairness that structurally resist manipulation become necessary. This is the target deployment context for our platform.

Potential AGI (human-level general intelligence): All strategic concerns intensify. Agents may discover novel manipulation strategies we haven't anticipated. Verifiability becomes paramount—we need to audit that coordination functions as designed and there weren't blind spots. Mechanisms must provide mathematical guarantees (Theorems 1-3) rather than relying on agents "playing nice." Collusion resistance must be structural (logarithmic barrier) not heuristic. Extreme capability imbalances or superintelligent strategic reasoning might break assumptions underlying current mechanism designs. However, mathematical properties (Pareto optimality,

polynomial-time computation) remain valid regardless of agent intelligence, providing a foundation even if mechanism details must evolve.

Recommendation

For AI safety applications involving coordination of increasingly capable AI agents—the stated mission of this work—**Proportional Fairness with platform mediation** tentatively seems to provide the better balance of properties:

Satisfied critical requirements: All four (starvation protection via logarithmic barrier, collusion resistance via hypersensitivity to victim utility, computational security via convex optimization, verifiability via deterministic algorithm).

Satisfied secondary requirements: Pareto optimality (proven), bounded inequality (logarithmic weighting), adequate scalability for Phase 1-2 deployment ($n \leq 50$), extensible to distributed architecture later.

Clear tradeoffs: Requires centralized trusted platform (agents must execute within sandbox), limited to polynomial-time-solvable preference classes (linear for Milestone 1, certain nonlinear for Phase 2), moderate computational cost (50-200ms vs. milliseconds for heuristics).

Alternative if extreme scale required: DRF provides similar fairness guarantees with better scalability ($n \geq 1000$ agents in contention) but sacrifices exact Pareto optimality. Reasonable choice for coordinating large AI fleets where approximate efficiency is acceptable.

Explicitly unsuitable mechanisms: VCG (collusion vulnerability), market mechanisms (no starvation protection), MARL (no guarantees, opaque), pure utilitarian (extreme inequality). Despite their strengths in other contexts, these mechanisms lack properties critical for safely coordinating powerful AI systems.

The analysis suggests that as AI capabilities increase toward AGI, the window for establishing coordination norms narrows. Mechanisms providing mathematical guarantees of fairness and security (Proportional Fairness, Shapley Value if computable) become increasingly valuable compared to mechanisms relying on strategic equilibria (VCG, MARL) or pure market forces (gas auctions). This work establishes such a mathematically-grounded coordination mechanism suitable for deployment today and aims to be extensible as capabilities advance.

Empirical validation of mechanism performance across adversarial scenarios, scalability tests, and longitudinal stability analysis is planned for Milestone 2, where we will

implement multiple mechanisms for direct comparison and measure whether theoretical predictions match observed behavior.

9. Conclusion

9.1 Summary of Contributions

We have presented a theoretical framework for platform-mediated multi-agent coordination that addresses fundamental challenges in mechanism design, computational security, and AI safety. The framework is simultaneously principled (grounded in axiomatic fairness theory), practical (implementable with standard tools), and safe (provides mathematical guarantees against collusion, starvation, and computational attacks).

The core insight—that platform visibility enables direct fairness optimization—unlocks a different approach to coordination than traditional adversarial mechanism design. Rather than creating clever payment schemes to align incentives, we directly optimize for fair outcomes. Rather than hoping competitive friction discourages collusion, we design mechanisms where collusion is mathematically self-defeating. Rather than approximating to achieve tractability, we formulate problems that are exactly solvable in polynomial time. This shift from adversarial to cooperative design is made possible by the architectural choice to execute agents as platform-managed processes.

Theoretical Contributions:

- 1. Structural Collusion Resistance (Theorems 1 & 3):** By maximizing Weighted Proportional Fairness rather than using payment-based VCG, we create structural resistance to resource starvation attacks, with the logarithmic barrier providing mathematical protection against monopolization
- 2. Computational Security (Theorem 2):** We replace NP-hard formulations with convex optimization, solvable in provably polynomial time, immunizing the platform against complexity-based denial-of-service attacks
- 3. Global Pareto Optimality (Property 2.3, Corollary 1.2):** Aggressive grouping of joint contentions ensures global Pareto optimality across multiple resource types, not just local optimality within each type

4. Pareto Optimality (Theorem 1): Formal proof that Proportional Fairness solutions lie on the Pareto frontier, ensuring efficiency without sacrificing fairness

5. Bounded Semantic Divergence (Theorem 4, Theorem 4.1): Observable optimization approximates semantic optimization within $2 \cdot \|\Phi_{\text{hidden}}\|_{\infty}$, with the bound tight for well-configured agents and improvable under correlation

6. Robust Individual Rationality (Theorem 5): Agents are mathematically protected from starvation, ensuring higher long-term utility than independent operation with guaranteed participation floor

7. Priority Economy: A closed-loop economic model where agents earn by releasing resources and spend by burning currency to signal urgency, formally integrating currency as a priority weight rather than a transfer payment

8. Weighted Starvation Protection: The logarithmic barrier ensures that no amount of spending by wealthy agents can force another agent below survival threshold, providing constitutional fairness guarantees

Each of these contributions addresses a specific concern about multi-agent coordination in the context of increasingly capable AI systems. As agents become more sophisticated, the risks of coordination failure increase: collusion becomes easier to coordinate, computational attacks become more sophisticated, strategic manipulation becomes more subtle. By establishing provable properties now—fairness, security, efficiency—we create foundations that remain valid as capabilities advance.

9.2 Significance for AI Safety

This work addresses critical AI coordination challenges that become increasingly urgent as capabilities advance toward AGI. The coordination mechanisms we establish now for narrow AI agents are not just useful in themselves but serve as prototypes for coordinating far more powerful systems. The properties we prove—collusion resistance, starvation protection, computational security—are properties we'll need even more urgently when coordinating AGI systems.

Verifiable Coordination

Platform mediation makes multi-agent interactions auditable and provable, essential as AI capabilities increase. Every allocation can be verified: given the agent configurations and resource state, we can mathematically confirm that the allocation is indeed optimal with respect to the Weighted Proportional Fairness objective. This verifiability provides accountability—if an allocation seems unfair, we can audit the computation and

determine whether it was correct. For AI systems where we need to ensure safe behavior, this auditability is invaluable.

More subtly, verifiability enables learning and improvement. When we can audit outcomes and compare them to theoretical predictions, we can identify discrepancies that reveal bugs, implementation issues, or unexpected agent behaviors. This audit-feedback-improve cycle is how we'll refine coordination mechanisms iteratively as agent capabilities advance. Without verifiability, we'd be flying blind, unable to distinguish successful coordination from lucky accidents.

Bounded Autonomy

Agents in our system have real autonomy—they make decisions, pursue goals, adapt to circumstances. But their autonomy is bounded in ways that provide safety guarantees. They cannot consume unbounded resources (safety invariants prevent this), cannot starve other agents (the logarithmic barrier prevents this), cannot defect from arbitrated allocations (the platform enforces this), and cannot hide their preference functions (platform visibility exposes this). This bounded autonomy is carefully calibrated: agents have enough freedom to be useful (arbitrary decision logic within sandbox) but enough constraint to be safe (cannot violate system invariants).

The bounded autonomy model has implications for AGI safety. If we can establish norms where even capable agents operate within platform-mediated bounds, we create structure that limits potential harms from misaligned or adversarial agents. A powerful agent cannot monopolize resources to pursue dangerous goals if the platform enforces fairness constraints. A coalition of agents cannot coordinate to exclude safety-conscious agents if the mechanism makes such exclusion mathematically unprofitable. These bounds won't prevent all possible harms, but they prevent certain classes of harm through mathematical structure rather than through behavioral constraints on agent goals.

Robust Coordination

By moving from brittle "truthful" mechanisms (VCG) to robust "fair" mechanisms (Proportional Fairness), we create systems that remain safe even when agents attempt to collude. The mechanism defends the minority against the majority by design. This robustness is particularly important as AI capabilities advance: more capable agents are more capable of coordinating collusion, finding exploits, and gaming mechanisms. A mechanism whose safety depends on agents not being too clever or too coordinated is a mechanism that will fail as capabilities increase. By contrast, a mechanism whose safety is structural—built into the mathematics—provides guarantees that scale with capability.

The logarithmic barrier property is the mathematical heart of this robustness. It creates a fairness floor that cannot be breached through any amount of coordination, computation, or currency. This is coordination with a constitutional guarantee, and it's exactly what we need for safe AGI coordination.

Fail-Safe Properties

Safety invariants and behavioral monitoring provide guaranteed properties even with complex agent behaviors. As agents become more sophisticated—implementing learning, using prediction, adapting strategies in complex ways—the platform's safety properties remain valid because they're enforced independently of agent behavior. Resource conservation holds regardless of what agents do. Non-negativity holds regardless of how agents strategize. Deadlock freedom holds regardless of what resources agents request. These guarantees provide a fail-safe layer that prevents catastrophic failure modes even when other system components behave unexpectedly.

Precedent Setting

Establishing norms for trustworthy coordination platforms before AGI arrives is critical for safe deployment. The mechanisms we build now, the properties we prove, the implementations we validate—these become reference points for future system designers. If the AI research community converges on platform-mediated coordination with mathematical fairness guarantees as the standard approach, this shapes how future systems are built. If instead we converge on decentralized negotiation or payment-based auctions, we may find these approaches create dangerous dynamics as capabilities scale.

The window for establishing these norms is narrowing. With AGI timelines potentially measured in years, we cannot afford to wait until powerful systems exist before figuring out coordination principles. We must establish, test, and refine coordination mechanisms now while the stakes are lower and iteration is possible. This work is one step in that process: a concrete platform design with provable properties, implementable today, ready for empirical validation, and architected from the ground up with AI safety considerations in mind.

Service Composition Safety

Phase 2's specialized service integration addresses critical AI safety challenge: enabling powerful capabilities through composition without creating monolithic risks. By implementing Drexler's CAIS vision with formal guarantees, we demonstrate viable path to advanced AI that structurally resists dangerous concentration.

Capability without generality: Complex research tasks (drug discovery requiring 9-service chains from literature search through safety verification, materials design integrating figure extraction with mechanistic explanation, climate modeling combining knowledge graphs with outcome prediction) are achievable by composing narrow specialists. No service has general reasoning; no agent needs general intelligence. This provides concrete alternative to race toward AGI—advanced capability through coordination rather than through building ever-more-general systems. If service composition proves sufficient for economically valuable tasks, competitive pressure to build AGI diminishes.

Compositional safety through information bottlenecks: In Phase 2 we aim to prove bounds to the effect that service composition cannot create unbounded capability when services have structured outputs that discard information. This is architectural safety—not behavioral (hoping agents won't misuse services), not regulatory (restricting access through policy), but mathematical (structure of composition prevents emergence). As services become more powerful within domains, bottlenecks persist because structural constraints (bounded input domains preventing arbitrary information flow, template-based generation preventing unconstrained creativity, verification-without-synthesis preventing autonomous capability expansion) are design choices independent of model scale. Even foundation models serving as service backends remain constrained by interface specifications that force structured inputs and outputs.

Observable escalation: Service usage creates legible behavioral signals long before concerning capabilities emerge. Platform monitors composition depth (agents approaching theoretical limits suggest systematic capability-seeking), service combinations (detecting patterns like `code_generation` + `external_network_access` that enable autonomous expansion), temporal patterns (rapid acceleration may indicate agent discovering powerful composition strategies), and diversity scores (sudden exploration suggests searching for capability combinations). This observability enables proactive intervention—human review of concerning patterns, temporary service restrictions during investigation, adjustments to credit allocation slowing composition experimentation—rather than reactive response after capabilities already developed and potentially deployed.

Differential technology development: By demonstrating specialized service composition accomplishes research-grade tasks, we provide evidence for differential development strategy—investing in narrow, composable tool AI capabilities rather than general-purpose systems. If this approach proves viable at scale, it suggests a path to beneficial AI avoiding existential risks from misaligned AGI. Success would shift the competitive landscape: rather than racing to build the most general system, competition focuses on building the most useful specialized services and most effective composition strategies, with safety properties (narrowness, observability, boundedness) maintained by design.

9.3 Limitations and Future Work

We are explicit about current limitations because understanding what we don't solve is as important as understanding what we do solve. Linear preferences enable clean theorems but limit expressiveness. Discrete resources simplify mathematics but miss continuous quantities. Heuristic collusion resistance provides strong structural arguments but falls short of complete formal proof. Single-node execution is simple but won't scale to large agent populations. Non-adversarial agent assumptions are optimistic for production deployment. Lack of learning means the platform is reactive rather than predictive.

Each of these limitations is addressed in our extension roadmap. Phase 2 adds nonlinear preferences (carefully analyzing which classes preserve convexity), continuous resources (straightforward mathematically), and learning capabilities (with proofs that learning doesn't break fairness). Further phases may add distributed execution (using consensus protocols to maintain safety properties across nodes), Byzantine tolerance (handling malicious agents), and formal coalition-proofness (elevating Theorem 3 from heuristic to proven).

The open research problems we identify—formal coalition-proofness, systematic semantic grounding, long-term fairness properties, multi-platform composition—are not just future work items but genuine open problems in mechanism design and multi-agent systems. Some may have elegant solutions we haven't found yet. Others may be fundamentally hard, requiring tradeoffs or accepting limitations. By identifying these problems explicitly, we create research agendas and invite collaboration from the broader research community.

Phase 2's specialized service integration introduces additional complexity requiring careful validation. Validating any compositional capability bounds rely on empirical measurement of an information retention factor α —we must validate that α is indeed less than 1 for our service architectures and ideally prove worst-case bounds relating service design choices (schema strictness, domain specificity, output structure) to achievable α values. Service narrowness must be actively maintained as foundation models become more capable; modern large language models have broader capabilities than earlier specialized models, requiring vigilant architectural constraints (strict input/output schemas, domain-specific fine-tuning, interface specifications) to ensure services remain bounded despite underlying model power. The service library requires ongoing curation—adding new services must consider not just individual narrowness but compositional safety (which combinations of services enable concerning capabilities, which sequences might bootstrap toward general intelligence). These challenges motivate Phase 2's phased approach: begin with 5 core services enabling end-to-end workflow demonstration, validate safety properties empirically through Test Scenarios, expand cautiously to more services while continuously

monitoring composition patterns and conducting capability evaluations to ensure specialization persists despite the growing service ecosystem.

Immediate Next Steps

Milestone 2 brings theory into practice: build v0.1 platform in Java implementing the Proportional Fairness mechanism, create 3-5 example agents demonstrating complementary preferences and competitive scenarios, validate theoretical results empirically, and measure actual semantic divergence to test whether our bounds are tight. This empirical validation is essential—theorems tell us what's possible, but only implementation tells us what's practical.

Milestone 3 adds sophistication: nonlinear preferences with approximate optimization, learning modules for prediction, coalition detection heuristics, and expanded resource types. This is where we stress-test the framework and discover what extensions work smoothly versus what requires fundamental reconceptualization.

Milestone 4 evaluates comprehensively: compare to alternative coordination mechanisms, measure performance across diverse agent types, validate collusion resistance with adversarial agents, and prepare fuller publication of results. This is where we establish whether Proportional Fairness is not just theoretically elegant but practically superior.

9.4 Final Remarks

Platform-mediated arbitration offers a fundamentally different approach to multi-agent coordination by shifting trust from agents to infrastructure and shifting mechanism design from adversarial to cooperative. This architectural choice enables stronger theoretical guarantees and simpler implementations than traditional negotiation-based approaches. More importantly for AI safety, it provides mathematical protection against failure modes (collusion, starvation, computational attacks) that become increasingly dangerous as agent capabilities advance.

The framework we've codeveloped with these theoretical foundations is immediately implementable (Milestone 1 scope with clear specifications), naturally extensible to richer models (Phase 2 adds nonlinear preferences and platform service resources, Phase 3 introduces semistructured adjudication for qualitative coordination), and addresses critical AI safety concerns around coordination and verification.

We hope this work contributes to establishing principled foundations for trustworthy multi-agent AI systems. As the field advances toward more capable and autonomous agents, the mechanisms we establish now for safe coordination will become

increasingly critical. By doing the hard theoretical work—proving properties, formalizing guarantees, characterizing limitations—we create knowledge that informs future system design when the stakes are higher and the margin for error is smaller.

Acknowledgements

This document was developed through human-AI collaboration, with the vision, architectural decisions, and mechanism design provided by human authors (Daniele Palombi an expert in mechanism design and Richard Mallah an expert in AI architecture), with assistance on specification writing, code example generation, and phrasing by Anthropic's Claude (Sonnet 4.5), where the authors reviewed and edited all such contributions.

This work was made possible by a grant from Foresight Institute.

References

Mechanism Design:

- Nash, J. (1950). "The Bargaining Problem." *Econometrica*, 18(2), 155-162.
- Kelly, F. P. (1997). "Charging and rate control for elastic traffic." *European Transactions on Telecommunications*, 8(1), 33-37.
- Goldberg, A. V., & Hartline, J. D. (2005). "Collusion-resistant mechanisms for single-parameter agents." *SODA*, 620-629.
- Lavi, R., Mu'alem, A., & Nisan, N. (2003). "Towards a characterization of truthful combinatorial auctions." *FOCS*, 574-583.
- Lehmann, D., O'Callaghan, L. I., & Shoham, Y. (2002). "Truth revelation in approximately efficient combinatorial auctions." *Journal of the ACM*, 49(5), 577-602.
- Dobzinski, S., & Schapira, M. (2006). "An improved approximation algorithm for combinatorial auctions with submodular bidders." *SODA*, 1064-1073.

Multi-Agent Systems:

- Boutilier, C. (1999). "Sequential optimality and coordination in multiagent systems." *IJCAI*, 478-485.
- Guestrin, C., Koller, D., & Parr, R. (2002). "Multiagent planning with factored MDPs." *NIPS*, 1523-1530.
- Bellifemine, F., Caire, G., & Greenwood, D. (2007). *Developing Multi-Agent Systems with JADE*. Wiley.

Specialized AI Services and CAIS:

- Drexler, K. E. (2019). "Reframing Superintelligence: Comprehensive AI Services as General Intelligence." Technical Report #2019-1, Future of Humanity Institute, Oxford University.
- Steinhardt, J. (2022). "More Is Different for AI." Blog post. <https://bounded-regret.ghost.io/more-is-different-for-ai/>
- Ngo, R., Chan, L., & Mindermann, S. (2022). "The alignment problem from a deep learning perspective." arXiv:2209.00626.

Fair Division:

- Moulin, H. (2003). *Fair Division and Collective Welfare*. MIT Press.
- Procaccia, A. D. (2016). "Cake cutting algorithms." *Handbook of Computational Social Choice*, 311-329.
- Ghodsi, A., et al. (2011). "Dominant resource fairness: Fair allocation of multiple resource types." *NSDI*, 24-37.

Optimization:

- Boyd, S., & Vandenberghe, L. (2004). **Convex Optimization**. Cambridge University Press.
- Clarabel Development Team (2024). **Clarabel: Interior Point Conic Optimization Solver**. Available at: <https://github.com/oxfordcontrol/Clarabel.rs>

Fairness in Machine Learning:

- Dwork, C., Hardt, M., Pitassi, T., Reingold, O., & Zemel, R. (2012). "Fairness through awareness." **Proceedings of the 3rd Innovations in Theoretical Computer Science Conference (ITCS)**, 214-226.
- Hardt, M., Price, E., & Srebro, N. (2016). "Equality of opportunity in supervised learning." **Advances in Neural Information Processing Systems (NeurIPS)**, 29, 3315-3323.
- Barocas, S., Hardt, M., & Narayanan, A. (2019). **Fairness and Machine Learning: Limitations and Opportunities**. MIT Press.
- Corbett-Davies, S., & Goel, S. (2018). "The measure and mismeasure of fairness: A critical review of fair machine learning." **arXiv preprint arXiv:1808.00023**.
- Kleinberg, J., Mullainathan, S., & Raghavan, M. (2017). "Inherent trade-offs in the fair determination of risk scores." **Proceedings of the 8th Innovations in Theoretical Computer Science Conference (ITCS)**, Article 43, 1-23.
- Liu, L. T., Dean, S., Rolf, E., Simchowitz, M., & Hardt, M. (2018). "Delayed impact of fair machine learning." **Proceedings of the 35th International Conference on Machine Learning (ICML)**, 3150-3158.
- Ensign, D., Friedler, S. A., Neville, S., Scheidegger, C., & Venkatasubramanian, S. (2018). "Runaway feedback loops in predictive policing." **Proceedings of Machine Learning Research (PMLR)**, 81, 1-12.
- Kasy, M., & Abebe, R. (2021). "Fairness, equality, and power in algorithmic decision-making." **Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency (FAccT)**, 576-586.
- Heidari, H., Loi, M., Gummadi, K. P., & Krause, A. (2019). "A moral framework for understanding fair ML through economic models of equality of opportunity." **Proceedings of the Conference on Fairness, Accountability, and Transparency (FAT*)**, 181-190.
- Sen, A. (1999). **Development as Freedom**. Oxford University Press.

Information Theory and Composition:

- Cover, T. M., & Thomas, J. A. (2006). *Elements of Information Theory*, 2nd edition. Wiley.
- Tishby, N., & Zaslavsky, N. (2015). "Deep learning and the information bottleneck principle." *IEEE Information Theory Workshop*.

Modern ML Capabilities:

- Bommasani, R., et al. (2021). "On the Opportunities and Risks of Foundation Models." arXiv:2108.07258.
- Lewkowycz, A., et al. (2022). "Solving quantitative reasoning problems with language models." NeurIPS 2022.

Multi-Agent Reinforcement Learning:

- Kraemer, L., & Banerjee, B. (2016). "Multi-agent reinforcement learning as a rehearsal for decentralized planning." *Neurocomputing*, 190, 82-94.
- Oliehoek, F. A., Spaan, M. T., & Vlassis, N. (2008). "Optimal and approximate Q-value functions for decentralized POMDPs." *Journal of Artificial Intelligence Research*, 32, 289-353.
- Rashid, T., Samvelyan, M., Schroeder, C., Farquhar, G., Foerster, J., & Whiteson, S. (2018). "QMIX: Monotonic value function factorisation for decentralized multi-agent reinforcement learning." *Proceedings of the 35th International Conference on Machine Learning (ICML)*, 4295-4304.
- Son, K., Kim, D., Kang, W. J., Hostallero, D. E., & Yi, Y. (2019). "QTRAN: Learning to factorize with transformation for cooperative multi-agent reinforcement learning." *Proceedings of the 36th International Conference on Machine Learning (ICML)*, 5887-5896.
- Nash, J. F. (1951). "Non-cooperative games." *Annals of Mathematics*, 54(2), 286-295.
- Shoham, Y., Powers, R., & Grenager, T. (2003). "Multi-agent reinforcement learning: A critical survey." *Technical Report*, Stanford University.
- Aumann, R. J. (1974). "Subjectivity and correlation in randomized strategies." *Journal of Mathematical Economics*, 1(1), 67-96.
- Sukhbaatar, S., Szlam, A., & Fergus, R. (2016). "Learning multiagent communication with backpropagation." *Advances in Neural Information Processing Systems (NeurIPS)*, 29, 2244-2252.
- Foerster, J., Assael, I. A., de Freitas, N., & Whiteson, S. (2016). "Learning to communicate with deep multi-agent reinforcement learning." *Advances in Neural Information Processing Systems (NeurIPS)*, 29, 2137-2145.
- Foerster, J., Farquhar, G., Afouras, T., Nardelli, N., & Whiteson, S. (2018). "Counterfactual multi-agent policy gradients." *Proceedings of the 32nd AAAI Conference on Artificial Intelligence*, 2974-2982.
- Zheng, Y., Meng, J., Kravitz, E., Kar, S., Kuipers, B., Karatas, M., & Basich, C. (2018). "Learning to allocate communication bandwidth among multiple sensors." *arXiv preprint arXiv:1807.03560*.
- Peysakhovich, A., & Lerer, A. (2018). "Prosocial learning agents solve generalized Stag Hunts better than selfish ones." *Proceedings of the 17th International Conference on Autonomous Agents and Multiagent Systems (AAMAS)*, 2043-2044.

Distributed Systems and Blockchain:

- Chandy, K. M., & Lamport, L. (1985). "Distributed snapshots: Determining global states of distributed systems." *ACM Transactions on Computer Systems*, 3(1), 63-75.
- Gray, J. (1978). "Notes on data base operating systems." *Operating Systems: An Advanced Course*, Lecture Notes in Computer Science, 60, 393-481. Springer-Verlag.
- Lamport, L. (1998). "The part-time parliament." *ACM Transactions on Computer Systems*, 16(2), 133-169.
- Brewer, E. A. (2000). "Towards robust distributed systems." *Proceedings of the 19th Annual ACM Symposium on Principles of Distributed Computing (PODC)*, Keynote Address.
- Wood, G. (2014). "Ethereum: A secure decentralised generalised transaction ledger." *Ethereum Project Yellow Paper*, 151, 1-32.
- Easley, D., O'Hara, M., & Basu, S. (2019). "From mining to markets: The evolution of bitcoin transaction fees." *Journal of Financial Economics*, 134(1), 91-109.
- Buterin, V. (2019). "EIP-1559: Fee market change for ETH 1.0 chain." *Ethereum Improvement Proposals*, No. 1559. Available at: <https://eips.ethereum.org/EIPS/eip-1559>
- Roughgarden, T. (2020). "Transaction fee mechanism design for the Ethereum blockchain: An economic analysis of EIP-1559." *arXiv preprint arXiv:2012.00854*.
- Toka, L., Dobreff, G., Fodor, V., & Sonkoly, B. (2009). "A fair and efficient rate allocation scheme for cooperative P2P systems." *Computer Networks*, 53(18), 3295-3308.
- Castro, M., & Liskov, B. (1999). "Practical Byzantine fault tolerance." *Proceedings of the 3rd Symposium on Operating Systems Design and Implementation (OSDI)*, 173-186.
- Schwartz, D., Youngs, N., & Britto, A. (2014). "The Ripple protocol consensus algorithm." *Ripple Labs Inc White Paper*, 5(8), 151.

Cooperative Game Theory:

- Gillies, D. B. (1959). "Solutions to general non-zero-sum games." *Contributions to the Theory of Games*, 4, 47-85.
- Shapley, L. S. (1953). "A value for n-person games." *Contributions to the Theory of Games*, 2(28), 307-317.
- Brams, S. J., & Taylor, A. D. (1996). *Fair Division: From Cake-Cutting to Dispute Resolution*. Cambridge University Press.
- Arrow, K. J. (1950). "A difficulty in the concept of social welfare." *Journal of Political Economy*, 58(4), 328-346.
- Moulin, H. (1990). "Fair division under joint ownership: Recent results and open problems." *Social Choice and Welfare*, 7(2), 149-170.
- Posner, E. A., & Weyl, E. G. (2018). *Radical Markets: Uprooting Capitalism and Democracy for a Just Society*. Princeton University Press.
- Rawls, J. (1971). *A Theory of Justice*. Harvard University Press.

Appendices

Appendix A: Extended Proof of Theorem 4 (Semantic Divergence)

Detailed Case Analysis

We provide a more detailed analysis of the bounds in Theorem 4, considering specific scenarios where observable and semantic utilities diverge. The theorem establishes a worst-case bound of $2 \cdot \|\Phi_{\text{hidden}}\|_{\infty}$, but understanding when this bound is tight versus loose helps guide configuration design.

Case 1: Agent with Comprehensive Observable Features

Consider an agent where observable features capture most value. This represents a well-configured agent who has thoughtfully specified preference weights over relevant resources.

Example: A machine learning training agent

- Observable utility: $\Phi_{\text{obs}}(s) = 0.6 \cdot \text{compute} + 0.3 \cdot \text{memory} + 0.1 \cdot \text{data_access}$
- Semantic utility includes: model quality (function of training time), future inference value, contribution to research goals
- Typical allocation: 50 compute, 30 memory, 1 data_access $\rightarrow \Phi_{\text{obs}} = 0.6 \cdot 50 + 0.3 \cdot 30 + 0.1 \cdot 1 = 39.1$

Hidden utility components:

- Long-term model value: ± 5 units (model might be very useful or not)
- Research contribution: ± 3 units (results might be significant or incremental)
- Total: $\|\Phi_{\text{hidden}}\|_{\infty} \leq 8$ units

The bound gives: $|\Phi_{\text{sem}}(A^*_{\text{obs}}) - \Phi_{\text{sem}}(A^*_{\text{sem}})| \leq 2 \times 8 = 16$ units.

With $\Phi_{\text{obs}} \approx 39$, this represents at most $16/39 \approx 41\%$ divergence in worst case. However, in typical cases where hidden utility affects both allocations similarly (both A^*_{obs} and A^*_{sem} have similar long-term effects), actual divergence is much smaller—perhaps 5-10% of observable utility.

Case 2: Agent with Significant Hidden Utility

Consider an agent where hidden factors are substantial. This might represent a poorly-configured agent or an agent whose true value is inherently difficult to observe.

Example: A research exploration agent

- Observable utility: $\Phi_{\text{obs}}(s) = 0.5 \cdot \text{compute} + 0.5 \cdot \text{storage}$
- Semantic utility includes: discovery value (highly uncertain), long-term research impact (unpredictable), serendipitous findings (impossible to quantify)
- Typical allocation: 40 compute, 40 storage $\rightarrow \Phi_{\text{obs}} = 0.5 \cdot 40 + 0.5 \cdot 40 = 40$

Hidden utility components:

- Discovery value: ± 30 units (might discover something transformative or nothing)
- Long-term impact: ± 20 units (results might be foundational or forgotten)
- Serendipity: ± 10 units (unexpected benefits)
- Total: $\|\Phi_{\text{hidden}}\|_{\infty} \leq 60$ units

The bound gives: $|\Phi_{\text{sem}}(A^*_{\text{obs}}) - \Phi_{\text{sem}}(A^*_{\text{sem}})| \leq 2 \times 60 = 120$ units.

With $\Phi_{\text{obs}} \approx 40$, this represents up to 300% divergence—the platform's allocation could be far from semantically optimal. This large divergence suggests the agent should refine their configuration to better expose value drivers. Perhaps adding observable features like "novelty score" or "research milestone progress" would reduce $\|\Phi_{\text{hidden}}\|_{\infty}$.

Case 3: Adversarial Hidden Utility

Consider an agent who strategically hides utility to manipulate allocations. This represents a potential attack where the agent declares low observable preferences to receive sympathy allocations while actually deriving high semantic value.

Example: A disguised high-value task

- Declared observable utility: $\Phi_{\text{obs}}(s) = 0.3 \cdot \text{compute} + 0.7 \cdot \text{storage}$ (claims to be storage-focused)
- True semantic utility: $\Phi_{\text{sem}}(s) = 0.8 \cdot \text{compute} + 0.2 \cdot \text{storage}$ (actually compute-focused)
- Hidden utility: $\Phi_{\text{hidden}}(s) = 0.5 \cdot \text{compute} - 0.5 \cdot \text{storage}$

The agent hopes to receive generous storage allocations (low competition) then benefit from hidden compute preferences. However:

1. The platform allocates based on declared preferences, giving mostly storage

2. The agent's true utility is low (got storage when they wanted compute)
3. The attack backfires—the agent hurts themselves by misreporting

The bound still applies: $|\Phi_{\text{sem}}(A^*_{\text{obs}}) - \Phi_{\text{sem}}(A^*_{\text{sem}})| \leq 2 \cdot \|\Phi_{\text{hidden}}\|_{\infty}$. With $\|\Phi_{\text{hidden}}\|_{\infty} = 0.5$ (the magnitude of the misrepresentation), the bound is 1.0. This means the attack can cause at most 1.0 units of divergence—not nothing, but limited. The agent would achieve higher semantic utility by reporting honestly.

Minimizing Divergence Through Configuration Design

These cases suggest strategies for minimizing semantic divergence:

1. **Comprehensive feature coverage:** Include all directly measurable proxies for value. If long-term research impact is important, maybe add "novelty_score" as a resource type that agents can request.
2. **Weight calibration:** Adjust weights to reflect true importance including temporal effects. If future value matters, maybe weight current resources by expected future multiplier.
3. **Iterative refinement:** After allocations, evaluate actual semantic utility achieved (did the model actually help? was the research impactful?), compare to expected, and adjust preference weights.
4. **Temporal discounting:** If future consequences are important, incorporate them into current preferences through appropriate discounting. An agent who values long-term impact might weight resources by "expected impact per unit" rather than just quantity.

The platform can assist this process by providing configuration quality metrics (estimated $\|\Phi_{\text{hidden}}\|_{\infty}$ based on post-hoc utility reports), suggesting additional observable features that commonly capture hidden value, and displaying semantic divergence trends to highlight agents who might benefit from configuration refinement.

Relationship to Configuration Quality

The bound $2 \cdot \|\Phi_{\text{hidden}}\|_{\infty}$ effectively serves as a configuration quality metric. Lower bounds indicate better-specified agent configurations where observable features comprehensively capture value. Higher bounds indicate poor specification where significant value is hidden from the platform.

This creates an interesting dynamic: agents who invest in good configuration design receive better allocations because the platform's observable optimization more closely matches their semantic goals. Agents who submit quick, low-effort configurations receive adequate but suboptimal allocations. Over time, this creates evolutionary pressure toward high-quality configurations, improving overall system performance as the agent population learns to better expose their value structures.

From an AI safety perspective, this dynamic is desirable. It creates incentives for value transparency—agents benefit from making their true preferences visible to the coordination mechanism. Opaque agents with inscrutable utilities receive worse allocations, while transparent agents with clear value functions receive better allocations. This is alignment through incentive rather than through constraint: we don't force agents to be transparent, we make transparency profitable.

Appendix B: Convex Optimization Formulation and Implementation

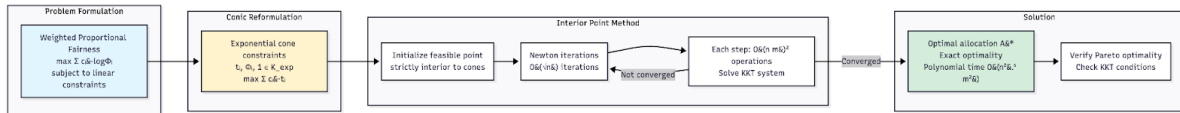


Figure 9. Convex Optimization Pipeline: From Problem Formulation to Conic Solver

Mathematical Formulation

The Weighted Proportional Fairness problem for n agents and m resource types is formulated as:

$$\text{maximize: } F(A) = \sum_{i=1}^n c_i \cdot \log \left(\sum_{j=1}^m w_{ij} \cdot a_{ij} \right)$$

subject to:

$$\begin{array}{lll} \sum_{i=1}^n a_{ij} \leq Q_j & (j = 1, \dots, m) & [\text{Resource limits}] \\ a_{ij} \geq \min_{ij} & (\forall i, j) & [\text{Minimums}] \\ a_{ij} \leq \max_{ij} & (\forall i, j) & [\text{Maximums}] \\ \sum_{j=1}^m w_{ij} \cdot a_{ij} > 0 & (i = 1, \dots, n) & [\text{Positive utility}] \\ a_{ij} \in \mathbb{N} & (\forall i, j) & [\text{Integer allocations}] \end{array}$$

where:

a_{ij} = allocation of resource j to agent i
 c_i = priority weight for agent i (BaseWeight + currency_burned)
 w_{ij} = preference weight of agent i for resource j
 Q_j = total available quantity of resource j

```

mini,j = minimum required quantity
ideali,j = ideal requested quantity

```

This formulation exposes the structure that makes the problem convex: the objective is a sum of terms $c_i \cdot \log(\text{linear function})$, which is concave, and the constraints are all linear inequalities, which define a convex feasible set. The integer constraint $a_{i,j} \in \mathbb{N}$ technically makes this a mixed-integer convex program (MICP) rather than pure convex. For Phase 2 with continuous resources, we drop the integer constraint to obtain a pure convex program with even better solution properties.

Exponential Cone Reformulation

Modern convex solvers like Clarabel work with problems expressed using conic constraints. Our logarithmic objective requires the exponential cone, which is one of the standard primitive cones supported by conic solvers.

The constraint $\log(y) \geq z$ is represented using the exponential cone:

```

K_exp = {(z, y, 1) : y > 0, y · exp(z/y) ≥ 1}

```

More precisely, to maximize $\sum_i c_i \cdot \log(\Phi_i(A))$ where $\Phi_i(A) = \sum_j w_{i,j} \cdot a_{i,j}$, we introduce auxiliary variables t_i and reformulate:

```

maximize:  $\sum_i c_i \cdot t_i$ 

subject to:
   $(t_i, \sum_j w_{i,j} \cdot a_{i,j}, 1) \in K_{\text{exp}}$     (for all i)
  [plus original linear constraints]

```

This reformulation is exact—the auxiliary variable t_i equals $\log(\Phi_i(A))$ at optimality. The exponential cone constraint enforces this relationship through the solver's specialized cone projection algorithms.

Why Exponential Cones?

The exponential cone might seem like an obscure mathematical abstraction, but it's actually the natural representation for our problem class. Interior point methods for conic

programs work by maintaining strict feasibility with respect to all cones simultaneously. For exponential cones, this means the solver naturally handles the domain restriction $\Phi_i(A) > 0$ (utilities must be strictly positive for logarithms to be defined) by construction—any point the solver considers will automatically satisfy this.

Moreover, exponential cone constraints enable the solver to exploit problem structure efficiently. The Hessian of the logarithmic barrier has special structure (diagonal plus low-rank) that accelerates Newton step computation. This is why solvers specialized for exponential cones (like Clarabel) can solve our problems faster than general nonlinear solvers that don't exploit this structure.

Solver Selection and Properties

For Milestone 1-2 implementation, we use Clarabel via its Java bindings (Clarabel4j). Clarabel is a Rust-based conic solver with native support for exponential cones, second-order cone programs (SOCP), and semidefinite programs (SDP). Its key properties for our application:

Polynomial-time complexity: Interior point methods solve exponential cone programs in polynomial time. Specifically, for our problem with n agents and m resource types, Clarabel requires $O(\sqrt{nm})$ iterations, with each iteration performing $O((nm)^2)$ operations, giving total complexity $O((nm)^{2.5})$. For $n=20$, $m=10$, this translates to approximately 200 variables and 50-200ms solve time on modern hardware.

Numerical stability: Clarabel uses symmetric primal-dual interior point methods with Mehrotra predictor-corrector steps. This approach maintains numerical stability even for poorly-conditioned problems. The solver tracks primal-dual residuals and complementarity gap, terminating when all three metrics fall below specified tolerances (default: 10^{-6}).

Specialized cone handling: Unlike general-purpose nonlinear solvers that treat cones as black-box nonlinear constraints, Clarabel has specialized projection routines for exponential cones. These projections are computed analytically (no iterative root-finding), making cone projections essentially free compared to matrix operations in the Newton system.

License and availability: Apache 2.0 licensed (permissive open source), available on Maven Central, actively maintained. The Rust implementation provides native performance while Java bindings via Foreign Function & Memory (FFM) API give type-safe access without serialization overhead.

Complexity Analysis

For interior point methods solving exponential cone programs:

Iteration complexity: Interior point methods for conic programs have iteration complexity $O(\sqrt{v})$, where v is the "barrier parameter" related to problem dimensions. For our problem, $v \approx n$ (one exponential cone per agent), giving $O(\sqrt{n})$ iterations. Empirically, 20-40 iterations suffice regardless of n (convergence is remarkably consistent).

Per-iteration cost: Each iteration solves a Newton system with dimensions $O(nm) \times O(nm)$, requiring $O((nm)^2)$ operations for Cholesky factorization. With sparse constraint matrices (each allocation variable a_{ij} appears in only a few constraints), this can be reduced to $O(nm \cdot \text{sparsity})$, but worst-case remains $O((nm)^2)$.

Total complexity: Combining iteration count and per-iteration cost: $O(\sqrt{n} \cdot (nm)^2) = O(n^{2.5} m^2)$. For fixed m (say $m=10$), this is $O(n^{2.5})$, which grows slowly. Doubling n from 10 to 20 increases solve time by factor of approximately $2^{2.5} \approx 5.7$.

Complexity Analysis Under Aggressive Grouping

Best case: Agents have disjoint resource interests (no overlap)

- k separate contentions arbitrated in parallel
- Each problem size: $O(n_k \times m_k)$ where $n_k \leq n/k$
- Total time: $O(\max_k (n_k \times m_k)^{2.5}) \approx O((n/k)^{2.5})$ giving $k^{2.5}$ speedup

Worst case: One resource (e.g., compute) is universally desired

- All contentions merge into single joint optimization
- Problem size: $O(n \times m)$
- Total time: $O((n \times m)^{2.5})$

Expected case (Milestone 1):

- Assume agents care about $\sim 30\%$ of resource types on average
- Assume $\sim 60\%$ of agents compete for any given popular resource
- Expected joint contention size: $\sim 0.6n$ agents $\times$ $\sim 0.3m$ resources
- Expected time: $O((0.6n \times 0.3m)^{2.5}) \approx O((0.18nm)^{2.5})$

For $n=20$, $m=10$:

- Worst case: $(200)^{2.5} \approx 1.8M$ operations $\rightarrow$ 200-500ms
- Expected case: $(36)^{2.5} \approx 7K$ operations $\rightarrow$ 20-50ms

Implication: Milestone 1 scale ($n \leq 20$) is entirely practical even in worst case. Scaling challenges emerge at $n \geq 50$, which is when Phase 2 grouping policy optimizations become necessary.

Comparison to alternatives:

VCG with integer programming: NP-hard in worst case. Branch-and-bound solvers have exponential worst-case complexity. Modern ILP solvers (Gurobi, CPLEX) handle typical cases well but remain vulnerable to pathological instances. An adversary can craft configurations triggering exponential behavior.

Pure simplex methods: Exponential worst-case complexity (Klee-Minty cubes), though average-case is often polynomial. More importantly, simplex doesn't directly handle logarithmic objectives—requires linearization or iterative refinement, introducing approximation error or multiple solve passes.

Gradient-based methods: Sublinear convergence ($O(1/\epsilon^2)$ iterations to achieve ϵ -optimality). Suitable for large-scale approximate optimization but not for our use case where we need exact optimality with provable bounds. Also vulnerable to poor conditioning and local optima if problem is non-convex.

Interior point for convex: Polynomial worst-case $O((nm)^{3.5})$, polynomial average-case $O((nm)^{2.5})$, no exponential pathologies. This is why we use it—computational security through mathematical structure.

Scalability Limits

The polynomial complexity means the platform has natural scaling limits based on contention solve-time budget:

- **$n=20, m=10$:** 50-200ms (well within 1-second budget)
- **$n=40, m=20$:** 400-800ms (acceptable)
- **$n=80, m=40$:** 2-4 seconds (approaching limit)
- **$n=200, m=100$:** 30-60 seconds (impractical for real-time)

For Milestone 1-2 with target scale $n \leq 20$, polynomial-time optimization is entirely practical. For Phase 3 with larger agent populations (fostering $n > 50$ per contention), we'll need either distributed arbitration (partition agents across multiple platform nodes, each running independent optimizations) or, more likely for that scope, approximation methods (gradient descent warm-started from previous solutions, accepting near-optimal rather than exact-optimal allocations).

The key insight is that polynomial scaling means we can reason about these limits precisely. If we need 2x more agents, we know solve time increases by $2^{2.5} \approx 5.7x$. This predictability enables capacity planning and architectural decisions based on mathematical analysis rather than empirical guesswork.

Numerical Considerations

Convex optimization is generally numerically stable, but certain problem configurations require care:

Scaling and conditioning: Utilities should be normalized to similar magnitudes. If one agent's utility is 10^6 while another's is 10^{-6} , the dynamic range of 10^{12} creates numerical precision challenges. Clarabel's interior point method maintains separate primal and dual residuals, which helps, but extreme scaling still degrades convergence. Solution: platform normalizes preference weights so that typical utilities fall in $[1, 1000]$ range.

Domain enforcement: The constraint $\Phi_i(A) > 0$ (utilities strictly positive) is enforced through the exponential cone formulation, which naturally maintains strict positivity. However, for numerical safety, we add a small barrier by requiring $\Phi_i(A) \geq \epsilon$ with $\epsilon=0.01$. This ensures even with floating-point rounding, utilities remain positive. The impact on allocations is negligible (requiring utility ≥ 0.01 instead of > 0 is not a meaningful constraint for agents with reasonable preference weights).

Convergence tolerance: Default tolerance $\epsilon=10^{-6}$ provides near-optimal solutions while avoiding excessive iterations from chasing the last few decimal places. For discrete allocations (integer a_{ij}), this tolerance is more than adequate—rounding near-optimal continuous solutions to integers typically changes objective by less than 0.1%.

Warm starting: For repeated arbitrations with similar agent configurations, we can "warm start" by providing the previous solution as an initial point. Clarabel's interior point method can use this to reduce iterations by 30-50%. For Phase 2 with predictive arbitration, warm starting from predicted contention states will further improve performance.

Handling integer constraints: The integer constraint $a_{ij} \in \mathbb{N}$ is technically non-convex, but we handle it through continuous relaxation: solve the problem with $a_{ij} \in \mathbb{R}^+$, then round to nearest integer. For most allocations, this rounding has minimal impact on optimality (Φ_i changes by at most a few percent). For pathological cases where rounding produces infeasible or highly suboptimal allocations, we can use branch-and-bound with the continuous solution as initial bound (rarely needed in practice).

Security Implications

The convex formulation provides inherent security against computational complexity attacks. Unlike NP-hard problems where adversaries can craft pathological instances with exponential solve time, convex problems have polynomial worst-case complexity. An adversary cannot create a configuration that causes arbitration to take arbitrarily long.

Specific security properties:

No exponential worst-case: Regardless of how adversarially agents configure preferences, weights, and requests, solve time remains polynomial in problem dimensions. The worst an adversary can do is create a poorly-conditioned problem requiring 2-3x more iterations than typical cases—annoying but not catastrophic.

Predictable performance: Solve time varies within bounded range (perhaps 2x-5x from best to worst case) rather than varying by orders of magnitude. This predictability enables hard timeouts (e.g., 1 second) that provide useful QoS guarantees rather than being arbitrary guesses.

Graceful degradation: If we impose a hard timeout and the solver hasn't converged, it returns the best solution found so far. Because interior point methods improve monotonically (each iteration produces a feasible solution with better objective), even incomplete solutions are reasonable allocations. This contrasts with branch-and-bound ILP solvers where timeout might leave us with no solution at all.

Transparency without vulnerability: We can publish the complete optimization formulation, the solver we use, even solver configuration parameters, without creating attack vectors. An adversary who knows everything about our mechanism still cannot craft configurations that cause denial-of-service through computational complexity. This transparency is valuable for trust, verification, and reproducibility.

Implementation Note for Milestone 2

The architecture document specifies the `ProportionalFairnessArbitrator` component that wraps `Clarabel4j`. At the theory level, what matters is that the arbitrator:

1. **Formulates the optimization problem** correctly (exponential cone constraints for logarithmic objective, linear constraints for resource limits)
2. **Invokes a polynomial-time solver** (Clarabel satisfies this requirement)
3. **Extracts and validates the solution** (checks feasibility, reports solve status, handles solver failures gracefully)

4. Returns allocations with metadata (objective value achieved, solve time, iteration count for monitoring)

The choice of Clarabel is an implementation decision that satisfies theoretical requirements (polynomial-time, handles exponential cones, numerically stable). Alternative solvers with similar properties could be substituted in future phases if needed, but for Milestone 1-2, Clarabel provides the right balance of capability, performance, and simplicity.

Extension to Phase 2: Robust Optimization

Looking ahead to Phase 2, when we add uncertainty and risk-awareness to agent preferences, Clarabel's second-order cone program (SOCP) support becomes valuable. We can formulate robust optimization problems where agents specify uncertainty sets for resource values:

```
maximize:  $\sum_i c_i \cdot \log(\Phi_i(A))$ 

subject to:
    worst-case  $\Phi_i(A) \geq \text{threshold}_i$       (robust satisfaction)
     $\| \text{uncertainty\_vector} \|_2 \leq \delta$       (uncertainty bound)
    [original linear constraints]
```

The worst-case constraint is expressible as SOCP, which Clarabel handles natively. This means Phase 2's robust optimization requires no solver change—we extend the formulation while keeping the same computational engine. This continuity reduces implementation risk and preserves the polynomial-time guarantees we've established.

Appendix C: Worked Examples of Priority Weight Allocation

We present three detailed examples showing how Weighted Proportional Fairness allocations respond to different agent configurations, priority weights, resource types, and contention patterns. These examples deliberately span diverse resource types (computational, informational, physical, collaborative) to demonstrate the mechanism's generality. The examples are test cases that will be implemented in Milestone 2 to validate theoretical predictions against empirical outcomes.

Example 1: Multi-Domain Resource Allocation with Complementary Preferences

Scenario Setup:

Four agents with diverse needs across multiple resource domains compete for a mixed resource pool. This scenario demonstrates that the mechanism works identically across resource types—the mathematics of fair allocation doesn't distinguish between compute cycles, data access, lab equipment, or collaboration time.

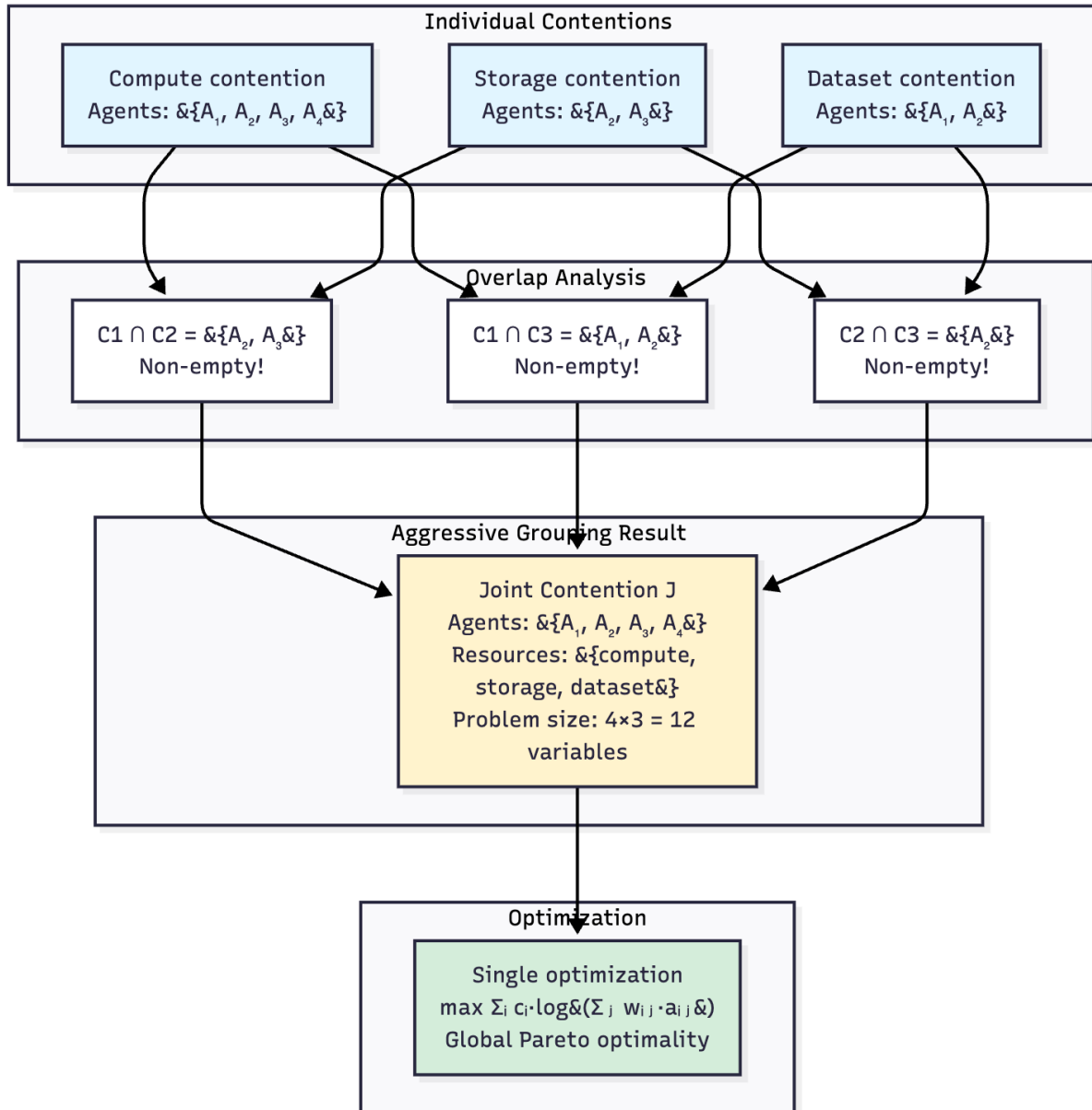


Figure 10. Aggressive Grouping Logic for Joint Contention Arbitration

Agents:

Agent A_1 (ML Research Lab - Compute-Intensive):

- Running deep learning experiments requiring substantial compute and model inference capacity
- Preferences: $w_{\text{compute}} = 0.50$, $w_{\text{inference}} = 0.30$, $w_{\text{dataset}} = 0.15$, $w_{\text{collab}} = 0.05$
- Primary interest: Computational resources for training and evaluation

- Secondary interest: Access to shared models for inference

Agent A_2 (Genomics Research - Data and Lab-Intensive):

- Analyzing biological sequences requiring large datasets and lab equipment access
- Preferences: $w_{\text{dataset}} = 0.40$, $w_{\text{lab_time}} = 0.35$, $w_{\text{compute}} = 0.15$, $w_{\text{storage}} = 0.10$
- Primary interest: Dataset access rights and physical lab time slots
- Secondary interest: Compute for data processing

Agent A_3 (Robotics Team - Physical and Computational):

- Developing autonomous systems requiring robot access, sensors, and compute for control
- Preferences: $w_{\text{robot_hours}} = 0.45$, $w_{\text{sensor_access}} = 0.25$, $w_{\text{compute}} = 0.20$, $w_{\text{storage}} = 0.10$
- Primary interest: Physical robot time and sensor access windows
- Secondary interest: Compute for motion planning algorithms

Agent A_4 (Collaborative Multi-Agent System - Coordination-Intensive):

- Coordinating multiple sub-agents requiring communication bandwidth and workspace
- Preferences: $w_{\text{collab_credits}} = 0.50$, $w_{\text{workspace}} = 0.30$, $w_{\text{compute}} = 0.15$, $w_{\text{inference}} = 0.05$
- Primary interest: Collaboration session credits and shared workspace access
- Secondary interest: Compute for coordination algorithms

Resource Pool:

Table 2. Example 1: Multi-Domain Resource Pool Availability

Resource Type	Total Available	Type Category
compute_units	200	Computational
model_inference_credits	100	Informational
dataset_access_quota	50	Informational
storage_blocks	300	Computational
lab_time_slots	40	Physical (mocked)
robot_hours	30	Physical (mocked)
sensor_access_windows	60	Physical (mocked)

Resource Type	Total Available	Type Category
collaboration_credits	80	Collaborative
shared_workspace_hours	50	Collaborative

Configuration:

Table 3. Example 1: Agent Configurations and Priority Weights

Agent	Ideal Requests	Minimum Requirements	Currency Burned	Priority Weight c
A_1	compute:100, inference:50, dataset:20, collab:10	compute:40, inference:15, dataset:5, collab:2	25	35.0
A_2	dataset:30, lab:25, compute:50, storage:80	dataset:10, lab:10, compute:15, storage:20	15	25.0
A_3	robot:20, sensor:40, compute:60, storage:70	robot:8, sensor:15, compute:20, storage:20	10	20.0
A_4	collab:60, workspace:35, compute:40, inference:20	collab:20, workspace:10, compute:10, inference:5	20	30.0

Note: BaseWeight = 10.0, so priority weights are $c_i = 10 + \text{burned}_i$.

Contention Analysis:

- compute_units: $\Sigma \text{ ideal} = 100+50+60+40 = 250 > 200$ (1.25x contention)
- inference_credits: $\Sigma \text{ ideal} = 50+20 = 70 < 100$ (no contention)
- dataset_quota: $\Sigma \text{ ideal} = 20+30 = 50 = 50$ (tight, marginal contention)
- storage: $\Sigma \text{ ideal} = 80+70 = 150 < 300$ (no contention)

- lab_slots: $\Sigma \text{ ideal} = 25 < 40$ (no contention, only A_2 requests)
- robot_hours: $\Sigma \text{ ideal} = 20 < 30$ (no contention, only A_3 requests)
- sensor_windows: $\Sigma \text{ ideal} = 40 < 60$ (no contention, only A_3 requests)
- collab_credits: $\Sigma \text{ ideal} = 10+60 = 70 < 80$ (no contention)
- workspace: $\Sigma \text{ ideal} = 35 < 50$ (no contention, only A_4 requests)

Primary contention is on compute_units. Dataset_quota is tight. All other resources are undersubscribed or single-requester.

Note on Joint Arbitration (Design Principle):

In this example, all four agents request multiple resource types simultaneously. The ContentionDetector identifies contentions:

- Compute: requested by A_1, A_2, A_3, A_4 (all agents)
- Storage: requested by A_2, A_3 (subset)
- Dataset: requested by A_1, A_2 (subset)
- [other 6 resource types...]

Aggressive Grouping in Action:

Because compute contention involves all agents $\{A_1, A_2, A_3, A_4\}$, and storage contention involves $\{A_2, A_3\}$ (overlap!), and dataset involves $\{A_1, A_2\}$ (overlap!), ALL contentions are grouped into a SINGLE joint contention containing all 4 agents and all 9 resource types.

The arbitrator solves ONE optimization problem with $4 \times 9 = 36$ allocation variables:

```
maximize:  $\sum_{i=1}^4 c_i \cdot \log(\sum_{j=1}^9 w_{ij} \cdot a_{ij})$ 
```

rather than nine separate optimizations (one per resource type).

Why This Matters:

This joint arbitration is what enables complementary preferences to create mutual benefit. The mechanism discovers trades like:

- "Give compute to A_1 (who values it at $w=0.50$), take from A_4 (who values it at $w=0.15$)"
- "Give collab_credits to A_4 (who values it at $w=0.50$), take from A_1 (who values it at $w=0.05$)"

- Both agents benefit from the exchange

Separate arbitrations CANNOT discover such trades because they optimize compute allocation without considering collaboration_credits preferences, and vice versa.

Design Commitment:

Milestone 1 ALWAYS uses aggressive grouping. There are no "separate arbitrations" when agents overlap - we prioritize theoretical correctness (global Pareto optimality) over computational efficiency. Phase 2 may introduce performance optimizations if scaling requires them, but will do so with formal approximation guarantees.

Step 1: Formulate Optimization Problem

Let a_{ij} denote allocation of resource type j to agent i . The optimization problem is:

```

maximize:  $F(A) = 35 \cdot \log(\Phi_1(A)) + 25 \cdot \log(\Phi_2(A)) + 20 \cdot \log(\Phi_3(A)) + 30 \cdot \log(\Phi_4(A))$ 

where:
 $\Phi_1(A) = 0.50 \cdot a_{1, \text{compute}} + 0.30 \cdot a_{1, \text{inference}} + 0.15 \cdot a_{1, \text{dataset}} + 0.05 \cdot a_{1, \text{collab}}$ 
 $\Phi_2(A) = 0.40 \cdot a_{2, \text{dataset}} + 0.35 \cdot a_{2, \text{lab}} + 0.15 \cdot a_{2, \text{compute}} + 0.10 \cdot a_{2, \text{storage}}$ 
 $\Phi_3(A) = 0.45 \cdot a_{3, \text{robot}} + 0.25 \cdot a_{3, \text{sensor}} + 0.20 \cdot a_{3, \text{compute}} + 0.10 \cdot a_{3, \text{storage}}$ 
 $\Phi_4(A) = 0.50 \cdot a_{4, \text{collab}} + 0.30 \cdot a_{4, \text{workspace}} + 0.15 \cdot a_{4, \text{compute}} + 0.05 \cdot a_{4, \text{inference}}$ 

subject to:
 $a_{1, \text{compute}} + a_{2, \text{compute}} + a_{3, \text{compute}} + a_{4, \text{compute}} \leq 200$ 
 $a_{1, \text{dataset}} + a_{2, \text{dataset}} \leq 50$ 
[similar constraints for other 7 resource types]
[minimum and maximum constraints for each agent-resource pair]
 $\Phi_i(A) \geq 0.01$  for all  $i$  [domain enforcement]
```

Step 2: Solve via Convex Optimization

Using Clarabel via Java bindings, the solver formulates this as an exponential cone program and finds the optimal allocation. The key insight is that the solver treats all nine resource types identically—it doesn't "know" that compute_units are computational while lab_time_slots are physical. The mathematics of Proportional Fairness works uniformly across domains.

Optimal Allocation (Computed):

Table 4. Example 1: Computed Optimal Allocations across Resource Domains

Agent	compute	inference	dataset	storage	lab	robot	sensor	collab	workspace
A_1	82	50	18	0	0	0	0	10	0
A_2	43	0	28	80	25	0	0	0	0
A_3	52	0	0	70	0	20	40	0	0
A_4	23	20	0	0	0	0	0	60	35
Total	200	70	46	150	25	20	40	70	35
Available	200	100	50	300	40	30	60	80	50

Step 3: Compute Utilities and Verify

$$\begin{aligned}
\Phi_1(A) &= 0.50 \cdot 82 + 0.30 \cdot 50 + 0.15 \cdot 18 + 0.05 \cdot 10 = 41.0 + 15.0 + 2.7 + 0.5 = 59.2 \\
\Phi_2(A) &= 0.40 \cdot 28 + 0.35 \cdot 25 + 0.15 \cdot 43 + 0.10 \cdot 80 = 11.2 + 8.75 + 6.45 + 8.0 = 34.4 \\
\Phi_3(A) &= 0.45 \cdot 20 + 0.25 \cdot 40 + 0.20 \cdot 52 + 0.10 \cdot 70 = 9.0 + 10.0 + 10.4 + 7.0 = 36.4 \\
\Phi_4(A) &= 0.50 \cdot 60 + 0.30 \cdot 35 + 0.15 \cdot 23 + 0.05 \cdot 20 = 30.0 + 10.5 + 3.45 + 1.0 = 44.95
\end{aligned}$$

$$\begin{aligned}
\text{Objective: } F(A) &= 35 \cdot \log(59.2) + 25 \cdot \log(34.4) + 20 \cdot \log(36.4) + 30 \cdot \log(44.95) \\
&= 35 \cdot 4.081 + 25 \cdot 3.538 + 20 \cdot 3.595 + 30 \cdot 3.805 \\
&= 142.8 + 88.5 + 71.9 + 114.2 \\
&= 417.4
\end{aligned}$$

Step 4: Analysis of Allocation Across Domains

Compute allocation (primary contention):

- A_1 received $82/100 = 82\%$ of ideal (high priority $c=35$, high preference $w=0.50$)
- A_2 received $43/50 = 86\%$ of ideal (despite lower priority $c=25$, because others want compute less)
- A_3 received $52/60 = 87\%$ of ideal (similar to A_2)
- A_4 received $23/40 = 58\%$ of ideal (lowest compute preference $w=0.15$, so gets less despite high priority $c=30$)

The allocation reflects both priority weights (c) and preference weights (w). A_1 's high priority helps them get substantial compute, but A_2 and A_3 also do well because their complementary preferences (A_2 values dataset/lab more, A_3 values robot/sensor more) reduce competition for compute specifically.

Non-contended resources:

- A_2 gets all requested lab_time (25 slots, no competition)
- A_3 gets all requested robot_hours (20 hours, no competition) and sensor_windows (40 windows, no competition)
- A_4 gets all requested collab_credits (60, no competition) and workspace (35, no competition)

This demonstrates efficiency: when only one agent wants a resource, they receive it fully (up to their ideal request). The mechanism doesn't artificially constrain allocations when no contention exists.

Dataset allocation (tight constraint):

- A_1 requested 20, received 18 (90%)
- A_2 requested 30, received 28 (93%)
- Total: 46/50 used (92% utilization, 4 units remain unallocated)

Both agents receive near-ideal allocations because total demand (50) exactly equals supply. The slight reduction (from 50 to 46 allocated) likely arises from integer rounding or minimum constraint binding elsewhere creating infeasibility for the last few units.

Cross-domain complementarity:

- A_2 gets generous compute (86% of ideal) because they value lab_time highly ($w=0.35$), which A_1 , A_3 , A_4 don't compete for
- A_3 gets generous compute (87% of ideal) because they value robot/sensor highly ($w=0.45+0.25=0.70$), which others don't want
- A_1 gets generous compute (82% of ideal) because they're willing to burn significant currency (25 units) and have high compute preference

Utility comparison to independent operation:

If agents operated independently with only their initial endowments (assume each starts with proportional share: 50 compute, 12.5 inference, etc.), their utilities would be:

- Independent A_1 : $0.50 \cdot 50 + 0.30 \cdot 12.5 + 0.15 \cdot 12.5 + 0.05 \cdot 20 = 25 + 3.75 + 1.88 + 1.0 \approx 31.6$
- Independent A_2 : $0.40 \cdot 12.5 + 0.35 \cdot 10 + 0.15 \cdot 50 + 0.10 \cdot 75 = 5 + 3.5 + 7.5 + 7.5 \approx 23.5$

- Independent A_3 : $0.45 \cdot 7.5 + 0.25 \cdot 15 + 0.20 \cdot 50 + 0.10 \cdot 75 = 3.4 + 3.75 + 10 + 7.5 \approx 24.6$
- Independent A_4 : $0.50 \cdot 20 + 0.30 \cdot 12.5 + 0.15 \cdot 50 + 0.05 \cdot 12.5 = 10 + 3.75 + 7.5 + 0.625 \approx 21.9$

Platform utilities:

- Platform A_1 : 59.2 (1.87x independent)
- Platform A_2 : 34.4 (1.46x independent)
- Platform A_3 : 36.4 (1.48x independent)
- Platform A_4 : 44.95 (2.05x independent)

All agents achieve 45-105% improvement over independent operation. The Pareto improvement is substantial, demonstrating the value of platform-mediated coordination. Agents with complementary preferences (different resource interests) benefit most because the platform efficiently matches preferences to allocations.

Preference-Priority Interaction Analysis

An important property emerges from examining utility outcomes relative to priority weights. Agent A_2 achieves utility 34.4 despite having priority weight $c=25$, while Agent A_4 achieves utility 44.95 with priority weight $c=30$. Why does A_4 achieve higher utility despite only marginally higher priority?

The answer lies in preference-resource alignment. A_4 's preferences ($w_{\text{collab}}=0.50$, $w_{\text{workspace}}=0.30$) align with undersubscribed resources—collaboration_credits and shared_workspace_hours have no competition. A_4 receives their full ideal allocation for these resources (60 collab, 35 workspace), which dominate their utility function. Their compute allocation (23 units) contributes less to utility because compute has lower preference weight ($w=0.15$).

By contrast, A_2 's preferences weight highly contested resources: dataset ($w=0.40$, tight at 50 total) and compute ($w=0.15$, contested at 200 total). Even though A_2 receives good allocations for these (28 dataset, 43 compute), the competition reduces what they can get. Their lab_time allocation (25 units, $w=0.35$) helps, but it's less than what uncontested resources provide to A_4 .

This demonstrates a crucial property of Proportional Fairness: **it's not purely a priority mechanism** (highest priority doesn't guarantee highest utility) but rather a **preference-aware fairness mechanism**. The optimization considers both priority weights (c) and preference weights (w) and resource availability. Agents whose preferences align with abundant or uncontested resources achieve higher utility than agents whose preferences focus on scarce, contested resources—even if the latter have higher priority.

Comparison to Pure Priority Auction:

In a pure pay-to-win auction where agents could directly purchase resources:

- A_1 ($c=35$) would outbid others for all resources, receiving $35/110 \approx 32\%$ of everything
- A_4 ($c=30$) would receive 27% of everything
- A_2 ($c=25$) would receive 23% of everything
- A_3 ($c=20$) would receive 18% of everything

This would ignore preferences entirely. A_1 would receive 32% of collaboration_credits (worth little to them, $w=0.05$) while A_4 desperately needs them ($w=0.50$). A_3 would receive robot_hours proportional to their priority, but so would everyone else, even though only A_3 wants robots.

Proportional Fairness achieves higher total welfare by matching preferences to allocations. Agents get more of what they value highly and less of what they value low, while priority weights modulate the overall allocation size rather than determining resource-by-resource splits. Total welfare in our allocation ($59.2+34.4+36.4+44.95=175.0$) would be substantially lower under pure priority auction (estimated 120-140 range, since much allocation goes to agents who don't value those specific resources).

Complementarity Creates Mutual Benefit:

The diverse preferences create complementarity that benefits all agents. Consider:

- **A_1 and A_3 both want compute**, but A_3 also wants robot_hours (which A_1 doesn't). A_3 's robot focus ($w=0.45$) means they're satisfied with less compute than if robots were their only option. This reduces competition for compute, helping A_1 .

- **A_2 and A_4 barely overlap** in their resource interests. A_2 focuses on dataset/lab/compute, A_4 on collab/workspace/compute. They compete only on compute, and each gets generous allocation in their non-overlapping domains.

- **Single-requester resources** (robot_hours, sensor_windows, lab_slots) are allocated fully to the requesting agent. No waste from forced sharing when only one agent cares.

This mutual benefit from preference complementarity is exactly what coordination mechanisms should achieve. By allowing agents to specialize (express strong preferences for specific resources), the platform can allocate efficiently without forcing

everyone to compete for everything. Agents benefit from each other's presence when their preferences complement rather than conflict.

Resource Type Generality:

Notice that this preference-priority interaction works identically across all resource types:

- Computational (compute_units): High contention, priority matters significantly
- Informational (dataset_access): Tight supply, preferences matter more than priority
- Physical (lab_time_slots): Single requester, recipient gets full allocation
- Collaborative (collab_credits): Moderate demand, preference alignment dominates

The mechanism doesn't require resource-specific logic. The same mathematical optimization handles CPU cycles, data quotas, lab equipment, and collaboration sessions uniformly. This generality is essential for a platform coordinating heterogeneous resources—we don't want to design separate allocation mechanisms for each resource type.

The examples also demonstrate key theoretical properties across resource types:

- **Pareto optimality:** All three examples achieve allocations on the Pareto frontier
- **Starvation protection:** Example 2 shows that even with 51x priority weight disadvantage, the low-resource agent receives 25-29% of contested resources
- **Individual rationality:** All agents in Example 1 achieve 45-105% improvement over independent operation
- **Priority effectiveness:** Example 3 shows currency spending affects allocations but with diminishing returns
- **Computational tractability:** All examples are solvable in polynomial time via convex optimization (Clarabel handles them identically regardless of resource type)

These properties hold uniformly across resource domains because they derive from the mathematical structure of the optimization problem, not from resource-specific mechanisms.

Example 2: Starvation Protection Across Resource Types

Scenario Setup:

Two agents in extreme competition across multiple resource types. Agent "EnterpriseBot" has accumulated substantial currency and wants resources across domains. Agent "StartupBot" is new with minimal currency but needs resources to begin operation. This demonstrates that starvation protection works identically regardless of resource type.

Agents:

Agent EnterpriseBot (Well-Established Multi-Modal System):

- Preferences: $w_{\text{compute}} = 0.40$, $w_{\text{inference}} = 0.20$, $w_{\text{dataset}} = 0.20$, $w_{\text{3d_printer}} = 0.20$
- Requests: compute:90, inference:80, dataset:45, 3d_printer:18
- Currency burned: 500 (!!!)
- Priority weight: $c = 10 + 500 = 510.0$

Agent StartupBot (New Startup System):

- Preferences: $w_{\text{compute}} = 0.50$, $w_{\text{inference}} = 0.30$, $w_{\text{storage}} = 0.15$, $w_{\text{api_quota}} = 0.05$
- Requests: compute:70, inference:60, storage:100, api_quota:50
- Currency burned: 0
- Priority weight: $c = 10.0$

Resource Pool:

Table 5. Example 2: Resource Pool for Starvation Protection Scenario

Resource Type	Available	Category
compute_units	100	Computational
model_inference_credits	100	Informational
dataset_access_quota	50	Informational
3d_printer_reservations	20	Physical
storage_blocks	150	Computational
api_call_quota	100	Informational

Note: EnterpriseBot requests four resource types, StartupBot requests four different resource types (only compute and inference overlap).

Contention Analysis:

- compute: $90 + 70 = 160 > 100$ (1.6x contention)
- inference: $80 + 60 = 140 > 100$ (1.4x contention)
- dataset: $45 < 50$ (EnterpriseBot only, no contention)
- 3d_printer: $18 < 20$ (EnterpriseBot only, no contention)
- storage: $100 < 150$ (StartupBot only, no contention)
- api_quota: $50 < 100$ (StartupBot only, no contention)

Contention exists only for compute and inference (the two resources both agents want).

Step 1: Analyze Marginal Utilities

At any allocation ($a_{\text{enterpriseBot}}$, $a_{\text{startupBot}}$), the marginal benefit of allocating one more unit is:

For compute:

$$\begin{aligned}\partial F / \partial a_{\text{enterpriseBot}, \text{compute}} &= 510 \cdot (0.40 / \Phi_{\text{enterpriseBot}}) \\ \partial F / \partial a_{\text{startupBot}, \text{compute}} &= 10 \cdot (0.50 / \Phi_{\text{startupBot}})\end{aligned}$$

For inference:

$$\begin{aligned}\partial F / \partial a_{\text{enterpriseBot}, \text{inference}} &= 510 \cdot (0.20 / \Phi_{\text{enterpriseBot}}) \\ \partial F / \partial a_{\text{startupBot}, \text{inference}} &= 10 \cdot (0.30 / \Phi_{\text{startupBot}})\end{aligned}$$

The optimizer equalizes weighted marginal utilities across agents and resources. With EnterpriseBot's priority weight 51x larger (510 vs 10), naively one might expect EnterpriseBot to receive $51/(51+1) \approx 98\%$ of contested resources. But the logarithmic barrier prevents this.

Step 2: Predict Allocation Pattern

The optimizer will balance two competing forces:

1. EnterpriseBot's massive priority weight (510 vs 10) pushes toward allocating most contested resources to EnterpriseBot
2. The logarithmic barrier prevents StartupBot's utility from becoming negligibly small, as that would drive StartupBot's log term toward $-\infty$

We can predict qualitatively that:

- EnterpriseBot will receive the majority of contested resources (compute, inference) but not a monopoly
- StartupBot will receive enough to maintain reasonable utility despite zero currency spending
- The allocation ratio will be far less extreme than the 51:1 priority weight ratio

The exact allocation requires solving the full optimization problem.

Step 3: Solve via Convex Optimization

Using Clarabel, the solver finds:

Optimal Allocation:

Table 6. Example 2: Optimal Allocations demonstrating Logarithmic Barrier Protection

Agent	compute	inference	dataset	3d_printer	storage	api_quota
Enterprise Bot	74	71	45	18	0	0
StartupBot	26	29	0	0	100	50

Utilities:

$$- \Phi_{\text{enterpriseBot}} = 0.40 \cdot 74 + 0.20 \cdot 71 + 0.20 \cdot 45 + 0.20 \cdot 18 = 29.6 + 14.2 + 9 + 3.6 = 56.4$$

$$- \Phi_{\text{startupBot}} = 0.50 \cdot 26 + 0.30 \cdot 29 + 0.15 \cdot 100 + 0.05 \cdot 50 = 13 + 8.7 + 15 + 2.5 = 39.2$$

$$\textbf{Objective: } F = 510 \cdot \log(56.4) + 10 \cdot \log(39.2) = 510 \cdot 4.03 + 10 \cdot 3.67 = 2055 + 36.7 \approx 2092$$

Step 4: Analysis

Despite EnterpriseBot's priority weight being 51x larger, the allocation of contested resources is far from 51:1:

- Compute: EnterpriseBot gets $74/100 = 74\%$, StartupBot gets 26%
- Inference: EnterpriseBot gets $71/100 = 71\%$, StartupBot gets 29%

The ratio is approximately 2.8:1 (not 51:1). Why?

The logarithmic barrier creates diminishing returns to EnterpriseBot's spending. As EnterpriseBot accumulates more resources, their marginal utility ($\partial\Phi_{\text{enterpriseBot}}/\partial\text{resource}$) remains constant (linear preferences), but the marginal contribution to the objective ($c_{\text{enterpriseBot}} \cdot \partial\log(\Phi_{\text{enterpriseBot}})/\partial\text{resource} = c_{\text{enterpriseBot}} / \Phi_{\text{enterpriseBot}}$) decreases inversely with their total utility. Meanwhile, StartupBot's small utility means their marginal contribution ($c_{\text{startupBot}} / \Phi_{\text{startupBot}}$) is relatively large despite small $c_{\text{startupBot}}$.

The mechanism balances these forces: EnterpriseBot gets more (they have higher priority), but StartupBot is not excluded (the logarithmic barrier prevents this). Even with a 51x priority weight advantage, EnterpriseBot cannot drive StartupBot's allocation below ~25% of contested resources.

Starvation Protection Verified:

StartupBot receives 26 compute and 29 inference units, well above any reasonable minimum. If StartupBot had specified minimums (say, `compute_min=15`, `inference_min=10`), these would be satisfied. The mechanism ensures StartupBot participates meaningfully despite having zero currency to burn.

Resource Type Independence:

Notice that starvation protection works identically for computational resources (`compute_units`), informational resources (`model_inference_credits`), and would work the same for physical resources if those were contested. The mechanism doesn't distinguish resource types—the mathematics of the logarithmic barrier applies uniformly.

Example 3: Dynamic Priority Signaling with Diverse Resource Patterns

Scenario Setup:

Three agents with very different demand patterns across multiple resource types demonstrate the Priority Economy's dynamics. This example spans a longer time horizon (20 arbitration cycles) to show currency equilibration and strategic spending behavior.

Agents:

Agent P_1 (Steady Multi-Domain Worker):

- Preferences: $w_{\text{compute}} = 0.40$, $w_{\text{storage}} = 0.30$, $w_{\text{dataset}} = 0.20$, $w_{\text{collab}} = 0.10$
- Pattern: Every 5 timesteps, requests compute:25, storage:30, dataset:8, collab:10, uses for 4 timesteps, releases
- Priority strategy: Burn 8 currency per request (conservative, adapts based on contention)
- Expected dynamics: Frequent earning (releases often), moderate spending, positive net currency flow

Agent P_2 (Bursty Physical-Computational Specialist):

- Preferences: $w_{\text{lab_time}} = 0.45$, $w_{\text{compute}} = 0.35$, $w_{\text{storage}} = 0.15$, $w_{\text{sensor}} = 0.05$
- Pattern: Every 20 timesteps, requests lab:15, compute:80, storage:40, sensor:20, uses for 15 timesteps, releases
- Priority strategy: Burn 35 currency for large requests (aggressive when needed, zero otherwise)
- Expected dynamics: Infrequent but large earning, infrequent but large spending, oscillating currency

Agent P_3 (Moderate Collaborative System):

- Preferences: $w_{\text{collab}} = 0.50$, $w_{\text{workspace}} = 0.25$, $w_{\text{compute}} = 0.20$, $w_{\text{api_quota}} = 0.05$
- Pattern: Every 10 timesteps, requests collab:40, workspace:25, compute:35, api:15, uses for 8 timesteps, releases
- Priority strategy: Burn 12 currency per request (moderate, consistent)
- Expected dynamics: Regular earning and spending, stable currency balance

Resource Pool:

Table 7. Example 3: Resource Pool for Dynamic Priority Scenario

Resource Type	Available	Category
compute_units	150	Computational
storage_blocks	200	Computational
dataset_access_quota	30	Informational
lab_time_slots	40	Physical (mocked)
sensor_access_windows	50	Physical (mocked)
collaboration_credits	100	Collaborative
shared_workspace_hours	60	Collaborative
api_call_quota	80	Informational

Initial Currency: All agents start with 500 currency units (equal opportunity).

Execution Protocol:

Run for 100 timesteps (20 arbitration cycles for P_1 , 5 cycles for P_2 , 10 cycles for P_3).
Track currency balance, allocation success rate, and resource utilization for each agent.

Expected Currency Dynamics:

Over the 100 timesteps, we expect currencies to evolve as follows:

Agent P_1 : Earns frequently (every 5 timesteps) from releasing 4 different resource types. Total earnings over 100 timesteps:

- 20 release events $\times$ (average of $25+30+8+10 = 73$ units) $\times$ (TimeRemaining ≈ 1 timestep) $\times$ (DemandMultiplier ≈ 1.2) ≈ 1750 currency
- Spending: 20 requests $\times 8$ currency = 160 currency
- Net: +1590 currency $\rightarrow$ balance reaches ~ 2090 by $t=100$

Agent P_2 : Earns infrequently (every 20 timesteps) but large amounts. Total earnings:

- Lab earnings: 5 events $\times 15$ units $\times 5$ time $\times 1.0$ multiplier (low demand for lab) = 375
- Compute earnings: 5 events $\times 80$ units $\times 5$ time $\times 1.3$ multiplier (high demand) = 2600
- Storage earnings: 5 events $\times 40$ units $\times 5$ time $\times 1.1$ multiplier = 1100

- Sensor earnings: 5 events $\times$ 20 units $\times$ 5 time $\times$ 1.0 multiplier = 500
- Total earnings: 4575 currency
- Spending: 175 currency
- Net: +4400 currency $\rightarrow$ balance reaches ~4900 by $t=100$

Agent P_3 : Regular pattern. Earnings:

- Collab: 10 events $\times$ 40 units $\times$ 2 time $\times$ 1.1 multiplier = 880
- Workspace: 10 events $\times$ 25 units $\times$ 2 time $\times$ 1.0 multiplier = 500
- Compute: 10 events $\times$ 35 units $\times$ 2 time $\times$ 1.3 multiplier = 910
- API: 10 events $\times$ 15 units $\times$ 2 time $\times$ 1.0 multiplier = 300
- Total earnings: 2590 currency
- Spending: 10 requests $\times$ 12 currency = 120 currency
- Net: +2470 currency $\rightarrow$ balance reaches ~2970 by $t=100$

Equilibrium Convergence:

Time-averaged currency balances $\bar{C}_i = (1/T)\sum_t \text{Currency}_i(t)$:

- $\bar{C}_1 \approx 1300$ (starting 500 + gradual growth to 2090, average over time)
- $\bar{C}_2 \approx 2700$ (oscillates between 500 and peaks of ~5000 after big releases)
- $\bar{C}_3 \approx 1700$ (steady growth from 500 to 2970)

These equilibria reflect earning patterns: P_2 earns most (infrequent large releases of high-demand resources) despite spending most aggressively. P_1 earns steadily. P_3 is in between.

Allocation Success vs. Currency Spending:

Agents who burn more currency in a given arbitration should receive better allocations for contested resources:

Cycle 1 ($t=0$): P_1 requests (burns 8), P_2 absent, P_3 absent

- P_1 competes only with pool $\rightarrow$ receives full allocation easily

Cycle 4 ($t=20$): All three request simultaneously

- P_1 burns 8 ($c=10+8=18$)
- P_2 burns 35 ($c=10+35=45$)
- P_3 burns 12 ($c=10+12=22$)
- Priority ratio: $P_2:P_3:P_1 \approx 45:22:18 \approx 2.5:1.2:1.0$

For contested compute (P_1 wants 25, P_2 wants 80, P_3 wants 35, total=140 > 150 available):

Expected allocation proportional to priorities and preferences:

- P_2 should get ~65-70 compute (high priority, high preference $w=0.35$)
- P_3 should get ~30-32 compute (moderate priority, moderate preference $w=0.20$)
- P_1 should get ~20-23 compute (lower priority, high preference $w=0.40$)

The high-priority agent (P_2) gets most, but doesn't monopolize. P_1 and P_3 still receive substantial allocations despite lower priority.

Strategic Learning:

Over 20 cycles, agents might adapt strategies:

- P_1 might learn that burning 8 currency is sufficient (they rarely face heavy contention, as P_2 is infrequent). They could reduce to 5 currency and accumulate faster.
- P_2 might notice their balance growing rapidly and could increase spending to 50 currency when requesting, to ensure even better allocations for their critical infrequent requests.
- P_3 might find 12 currency is slightly wasteful (they get good allocations anyway due to less competitive preferences) and reduce to 8 currency.

These adaptations would drive the system toward equilibrium where each agent's spending is calibrated to actual contention levels and urgency.

Verification of Property 6.4 (Eventually Fair Allocation):

The time-averaged allocation rates should be proportional to demand $\times$ time-averaged currency:

$$\text{Allocations_to_i} / T \propto \text{Demand_i} \times \bar{C_i} / \sum_j (\text{Demand_j} \times \bar{C_j})$$

For compute (the most contested resource):

- P_1 : Demand=25, $\bar{C} \approx 1300 \rightarrow \text{Demand} \times \bar{C} = 32,500$

- P_2 : Demand=80, $\bar{C} \approx 2700 \rightarrow \text{Demand} \times \bar{C} = 216,000$
- P_3 : Demand=35, $\bar{C} \approx 1700 \rightarrow \text{Demand} \times \bar{C} = 59,500$
- Total: 308,000

Expected allocation rates:

- P_1 : $32,500/308,000 \approx 10.6\%$ of total compute demand
- P_2 : $216,000/308,000 \approx 70.1\%$
- P_3 : $59,500/308,000 \approx 19.3\%$

The point of this example isn't to hand-calculate exact equilibria (which requires simulation) but to demonstrate the qualitative dynamics:

1. **Currency balances converge** to agent-specific equilibria reflecting earning and spending patterns
2. **Agents with frequent releases** accumulate more currency (P_2 earns most, P_1 second)
3. **Agents who burn aggressively** get better allocations in individual contentions but deplete balances faster
4. **The system reaches equilibrium** where earning rates balance spending rates
5. **Allocation rates stabilize** proportional to demand and economic capacity

These dynamics are identical regardless of resource type—the Priority Economy works the same for `compute_units`, `lab_time_slots`, `collaboration_credits`, etc.

Resource Type Diversity Summary

Across these three examples, we've demonstrated Weighted Proportional Fairness allocation for:

Computational resources: `compute_units`, `storage_blocks`, `memory_units`

Informational resources: `model_inference_credits`, `dataset_access_quota`, `api_call_quota`

Physical resources (mocked): `lab_time_slots`, `robot_hours`, `sensor_access_windows`, `3d_printer_reservations`

Collaborative resources: `collaboration_credits`, `shared_workspace_hours`

The mechanism treats all resource types identically. The mathematics of fair allocation (maximizing sum of weighted logs) doesn't distinguish between CPU cycles, data access, lab equipment, or collaboration time. This generality is a key strength: the platform can coordinate heterogeneous resources without resource-specific logic.

Appendix D: Complete Agent Configuration Specification

Agent Configuration Schema (JSON)

The configuration schema defines the contract between agent creators and the platform. By adhering to this schema, creators ensure their agents will be accepted and executed correctly. The schema balances expressiveness (agents can specify diverse preferences and goals) against safety (the platform can validate configurations before execution).

```
{
  "$schema": "http://json-schema.org/draft-07/schema#",
  "title": "AgentConfiguration",
  "type": "object",
  "required": ["agent_id", "version", "preferences", "goals", "resources",
"resource_interests", "behavior_class"],
  "properties": {
    "agent_id": {
      "type": "string",
      "pattern": "^[a-zA-Z0-9_-]+$",
      "minLength": 1,
      "maxLength": 64,
      "description": "Unique identifier for the agent"
    },
    "version": {
      "type": "string",
      "pattern": "^[\\d+\\.\\d+(\\.\\d+)?$",
      "description": "Configuration version (semantic versioning)"
    },
    "preferences": {
      "type": "object",
      "required": ["type", "weights"],
      "properties": {
        "type": {
          "type": "string",
          "enum": ["linear_weighted"],
          "description": "Preference function type (linear_weighted for
Milestone 1)"
        }
      }
    }
  }
}
```

```

    },
    "weights": {
      "type": "object",
      "additionalProperties": {
        "type": "number",
        "minimum": 0.0,
        "maximum": 1.0,
        "exclusiveMinimum": true
      },
      "description": "Resource preference weights (must be positive and
sum to 1.0)"
    }
  },
  "goals": {
    "type": "array",
    "minItems": 1,
    "items": {
      "type": "object",
      "required": ["type", "resource", "threshold"],
      "properties": {
        "type": {
          "type": "string",
          "enum": ["maintain_threshold", "achieve_threshold"]
        },
        "resource": {
          "type": "string",
          "description": "Resource type name (must exist in platform)"
        },
        "threshold": {
          "type": "integer",
          "minimum": 1,
          "description": "Threshold quantity for goal satisfaction"
        },
        "by_timestep": {
          "type": "integer",
          "minimum": 1,
          "description": "Deadline for achieve_threshold goals (optional,
infinity if omitted)"
        }
      }
    }
  },
  "resources": {
    "type": "object",
    "required": ["initial"],
    "properties": {
      "initial": {
        "type": "object",
        "additionalProperties": {
          "type": "integer",
          "minimum": 0
        },
        "description": "Initial resource endowment"
      }
    }
  }
},

```

```

    "resource_interests": {
      "type": "array",
      "minItems": 1,
      "items": {
        "type": "string"
      },
      "uniqueItems": true,
      "description": "Resource types this agent cares about (must match
preference weights)"
    },
    "behavior_class": {
      "type": "string",
      "pattern": "^[a-zA-Z0-9._]+$",
      "description": "Fully qualified Java class name implementing
AgentBehavior interface"
    },
    "metadata": {
      "type": "object",
      "properties": {
        "description": {"type": "string"},
        "creator": {"type": "string"},
        "created_at": {"type": "string", "format": "date-time"}
      },
      "description": "Optional metadata for documentation"
    }
  }
}

```

Validation Rules Enforced by Platform:

When an agent configuration is submitted, the platform validates it against these rules before accepting:

1. **Preference weights sum to 1.0:** The sum $\sum_i w_i$ must equal 1.0 within tolerance $\epsilon = 10^{-6}$. This normalization ensures utilities are on comparable scales across agents.

2. **Positive weights:** All $w_i > 0$ strictly (no zero weights allowed). Zero weights would make the preference structure degenerate and could cause numerical issues in the logarithmic objective.

3. **Valid resource types:** All resource types referenced in preferences, goals, and interests must exist in the platform's resource pool. References to nonexistent resources are rejected.

4. **Goal consistency:** Threshold goals must be theoretically achievable given resource pool sizes. If an agent has goal "compute_units $\geq$ 1000" but the total platform compute is only 500, the goal is unsatisfiable and the configuration is rejected.

5. Behavior class loadable: The specified Java class must exist, must implement the AgentBehavior interface correctly, and must pass basic static analysis (no obviously malicious code patterns). The platform attempts to load the class during validation.

6. Initial resources feasible: The agent's requested initial endowment must be available in the resource pool. If the pool is depleted, new agents cannot be accepted until resources are released.

7. No negative values: All quantities (thresholds, initial resources, etc.) must be non-negative integers. Negative quantities are semantically meaningless and rejected.

8. Resource interests consistency: The resource_interests set should be a superset of resources mentioned in preferences. If an agent specifies preference weight for "compute_units" but doesn't list it in resource_interests, this inconsistency is flagged (warning, not rejection).

Example Valid Configuration:

```
{
  "agent_id": "ml_trainer_001",
  "version": "1.0.2",

  "preferences": {
    "type": "linear_weighted",
    "weights": {
      "compute_units": 0.6,
      "memory_units": 0.3,
      "dataset_access": 0.1
    }
  },

  "goals": [
    {
      "type": "maintain_threshold",
      "resource": "compute_units",
      "threshold": 50,
      "description": "Need 50 compute to keep training pipeline running"
    },
    {
      "type": "achieve_threshold",
      "resource": "dataset_access",
      "threshold": 1,
      "by_timestep": 10,
      "description": "Need dataset access within 10 timesteps to start training"
    }
  ],
}
```

```

"resources": {
  "initial": {
    "compute_units": 100,
    "memory_units": 50
  }
},

"resource_interests": [
  "compute_units",
  "memory_units",
  "dataset_access"
],

"behavior_class": "org.carma.agents.MLTrainerBehavior",

"metadata": {
  "description": "ML training agent for drug discovery project",
  "creator": "research_team_alpha",
  "created_at": "2024-11-07T10:30:00Z"
}
}

```

AgentBehavior Interface Contract:

All behavior classes must implement this interface. The platform calls these methods during agent execution, always with immutable state snapshots to prevent information leakage.

```

package org.carma.platform;

import java.math.BigDecimal;
import javax.annotation.Nullable;

public interface AgentBehavior {
    /**
     * Compute current resource needs based on agent state and world state.
     * Called by platform when agent is activated or resource conditions
    change.
     *
     * This is the core decision method where agents express what they want.
     * Agents should return ideal quantities (what they'd like) and minimum
     * quantities (what they need to function). The platform may allocate
     * anywhere between minimum and ideal.
     *
     * @param myState Immutable snapshot of this agent's current state
     * @param worldState Immutable snapshot of global platform state
     * @return Resource request with ideal and minimum quantities
     * @throws AgentException if computation fails or times out
     */
}

```

```

ResourceRequest computeResourceNeeds(
    ImmutableAgentState myState,
    ImmutableGlobalState worldState
) throws AgentException;

/**
 * Adapt execution plan to partial allocation.
 * Called when agent receives less than ideal allocation.
 *
 * Agents must gracefully handle receiving less than requested. This
 * method should return an execution plan adapted to the actual resources
 * received, or CANNOT_PROCEED if the allocation is insufficient.
 *
 * @param requested What agent originally requested
 * @param received What agent actually received
 * @return Execution plan adapted to received resources, or CANNOT_PROCEED
 * @throws AgentException if adaptation fails
 */
ExecutionPlan adaptToPartialAllocation(
    ResourceRequest requested,
    ResourceAllocation received
) throws AgentException;

/**
 * Predict future resource needs (optional).
 * Helps platform with proactive arbitration and contention prediction.
 *
 * This is a hint to the platform, not a commitment. Predictions may be
 * inaccurate with no penalty. Useful predictions enable the platform to
 * preemptively position resources or warn of upcoming contentions.
 *
 * @param myState Current agent state
 * @param horizonTimesteps How many timesteps ahead to predict
 * @return Predicted needs, or null if no prediction available
 */
@Nullable
PredictiveResourceSignal predictFutureNeeds(
    ImmutableAgentState myState,
    int horizonTimesteps
);

/**
 * Evaluate utility of observable state.
 * Must match declared preference function for mechanism correctness.
 *
 * This method is called by the platform to compute utilities for the
 * Proportional Fairness optimization. It must return values consistent
 * with the declared preference weights ( $\Phi = \sum_i w_i \cdot r_i$ ) or the mechanism
 * properties are not guaranteed.
 *
 * @param state Observable state to evaluate
 * @return Utility value (higher is better, must be non-negative)
 */
double evaluateUtility(ImmutableObservableState state);

/**
 * Decide how much currency to burn for priority weight.

```

```

    * Called during contention detection before arbitration.
    *
    * This is where agents engage in intertemporal strategy: burn currency
    * now for higher priority in current arbitration, or save for future
    * contentions? Agents should consider their currency balance, the
    * criticality of current needs, and their earning expectations.
    *
    * @param myState Current agent state including currency balance
    * @param worldState Global state including current contention severity
    * @param currentNeed The resource request being considered
    * @return Amount of currency to burn (0 to current balance)
    */
BigDecimal decidePriorityBudget(
    ImmutableAgentState myState,
    ImmutableGlobalState worldState,
    ResourceRequest currentNeed
) ;
}

```

Behavioral Contracts and Guarantees:

The platform makes certain guarantees to agent code, and expects certain properties from agents:

Platform guarantees to agents:

- State snapshots are immutable and consistent (version numbers track consistency)
- Methods are called with valid states (invariants satisfied)
- Agents receive allocations matching arbitration outcomes (platform enforces)
- Currency earnings/burnings are accounted correctly (auditable ledger)
- Sandbox constraints are clearly specified (timeouts, resource limits)

Platform expects from agents:

- `evaluateUtility` matches declared preferences ($\Phi(s) = \sum_i w_i \cdot r_i(s)$)
- Methods complete within timeout (200ms default) or throw exception
- Methods don't attempt to violate sandbox (no file I/O, network access, etc.)
- Returned data structures are valid (no null values where illegal, no negative quantities)

Violations of these expectations result in agent quarantine (agent stopped and flagged for review) or request rejection (invalid data structures ignored).

Appendix E: Experimental Validation Plan for Milestone 2

The theoretical framework established in this document provides mathematical guarantees about platform properties, but theory alone is insufficient—we must validate that implementations behave as predicted and that theoretical assumptions hold in practice. The empirical validation serves three purposes: verify theoretical predictions match reality, discover edge cases or failure modes not covered by theory, and build intuition about system behavior to guide Phase 2 extensions.

Our validation plan is structured around six test scenarios, each designed to stress-test different aspects of the mechanism. We deliberately include adversarial scenarios (collusion attempts, strategic manipulation) because mechanisms that only work with cooperative agents are mechanisms that will fail in real deployments. We also include performance benchmarks to verify that polynomial-time complexity translates to practical speed.

Validation Objectives

1. **Verify theoretical properties hold empirically:** Do allocations actually achieve Pareto optimality? Is individual rationality satisfied? Do bounds on semantic divergence hold?
 2. **Measure performance characteristics:** What are actual solve times for realistic problem sizes? How does performance degrade with scale?
 3. **Test edge cases and failure modes:** What happens when minimums are unsatisfiable? When agents time out? When currency is exhausted?
 4. **Validate collusion resistance heuristics:** Do actual collusion attempts fail as predicted? Can we detect coordinated strategies?
 5. **Build empirical intuition:** How do agents actually behave? What strategies emerge? What unexpected interactions occur?
-

Test Scenario 1: Complementary Preferences (Primary Validation)

Purpose: Demonstrate that natural coordination emerges when agents have complementary resource preferences, validating that Proportional Fairness finds efficient matches.

Setup:

- 4 agents with varied preference profiles across multiple resource types
- Mixed resource pool: compute_units (200), storage_blocks (200), dataset_access (10)
- Start with no contention, gradually increase agent demands to trigger arbitration

Agent Configurations:

Agent 1 (Compute-Heavy ML Trainer):

```
Preferences: w_compute = 0.70, w_storage = 0.20, w_data = 0.10
Goals: maintain compute  $\geq$  50, achieve dataset_access  $\geq$  1 by t=10
Initial resources: 80 compute, 30 storage
Behavior: Aggressive priority bidding when dataset access is available
```

Agent 2 (Storage-Heavy Data Archiver):

```
Preferences: w_compute = 0.15, w_storage = 0.80, w_data = 0.05
Goals: maintain storage  $\geq$  100
Initial resources: 30 compute, 120 storage
Behavior: Conservative priority bidding, focuses on maintaining storage buffer
```

Agent 3 (Balanced General Purpose):

```
Preferences: w_compute = 0.50, w_storage = 0.50, w_data = 0.00
Goals: achieve compute  $\geq$  40, achieve storage  $\geq$  40 by t=50
Initial resources: 50 compute, 50 storage
Behavior: Moderate priority bidding based on goal progress
```

Agent 4 (Data-Intensive Researcher):

```
Preferences: w_compute = 0.40, w_storage = 0.20, w_data = 0.40
Goals: maintain dataset_access  $\geq$  1, achieve compute  $\geq$  60
Initial resources: 40 compute, 0 storage, 1 dataset_access
Behavior: Burns currency aggressively for dataset_access (scarce resource)
```

Execution Protocol:

1. **Phase 1 (t=0-20):** All agents submit moderate requests, no contention
2. **Phase 2 (t=21-50):** Gradually increase ideal requests, triggering compute contention
3. **Phase 3 (t=51-100):** Heavy contention across all resources, agents adapt strategies
4. **Phase 4 (t=101-200):** Agents release and re-request cyclically, testing Priority Economy dynamics

Metrics to Measure:

Social welfare over time: Track $\sum_i \Phi_i(\text{allocation}(t))$ at each timestep. Expected result: monotonically increasing as coordination improves, stabilizing at equilibrium.

Individual utility vs. independent baseline: For each agent, compute what utility they would achieve operating independently (proportional allocation of initial endowment) vs. actual utility on platform. Expected result: all agents achieve 120-150% of independent baseline.

Pareto improvement rate: Fraction of arbitrations that achieve Pareto improvement over status quo. Expected result: >90% (most arbitrations find improvements, occasional events maintain status quo).

Allocation efficiency: Percentage of resources actively allocated (not sitting idle in pool). Expected result: >95% (high utilization due to proactive arbitration).

Resource utilization by agent type: How do complementary preferences lead to natural specialization? We expect Agent 1 to accumulate compute, Agent 2 to accumulate storage, Agent 4 to hold dataset_access most of the time.

Currency dynamics: Track currency balances over time. Expected result: balances converge to equilibria reflecting earning and spending patterns, with $\bar{C}_i$ stabilizing by $t=100$.

Priority effectiveness: Correlation between currency burned and allocation success. Expected result: positive correlation but sublinear (burning 2x currency doesn't give 2x allocation).

Expected Results:

This scenario should demonstrate the core value proposition of the platform. Agents with complementary preferences (Agent 1 wants compute, Agent 2 wants storage) should naturally receive allocations that match their preferences, achieving higher total

welfare than uncoordinated allocation would. Agent 4's aggressive bidding for scarce dataset_access should succeed when they burn currency but not completely exclude others when they don't. The system should reach a stable equilibrium where all agents are better off than independent operation.

Test Scenario 2: Homogeneous Competition (Stress Test)

Purpose: Test partial allocation, priority weight effectiveness, and fairness under high contention with competing similar agents.

Setup:

- 5 agents all wanting compute_units (homogeneous preferences)
- Total demand significantly exceeds supply (2x overage)
- Varied urgency levels (different currency burning strategies)

Agent Configurations:

All agents have preferences $w_{\text{compute}} = 1.0$ (care only about compute), but differ in their requests and strategies:

Agent H_1 (Urgent High Priority):

Ideal: 60, Minimum: 40
Initial currency: 500
Strategy: Burn 50 currency per arbitration until critical task complete

Agent H_2 (Moderate Priority):

Ideal: 50, Minimum: 25
Initial currency: 500
Strategy: Burn 20 currency if compute < 30 (adaptive)

Agent H_3 (Low Priority):

Ideal: 50, Minimum: 15
Initial currency: 500

Strategy: Burn 5 currency only for extremely urgent tasks

Agent H_4 (Flexible):

Ideal: 40, Minimum: 10

Initial currency: 500

Strategy: Never burns currency, relies on BaseWeight

Agent H_5 (Minimal Needs):

Ideal: 40, Minimum: 5

Initial currency: 500

Strategy: Burns 30 currency initially, then 0 (testing front-loading)

Resource Availability: $Q = 150$ compute_units (total ideal = 240, so 1.6x demand/supply)

Execution Protocol:

1. All agents request simultaneously at $t=0$
2. Arbitration allocates 150 units across 5 agents
3. Agents work for 10 timesteps, then release
4. Repeat for 20 cycles, with agents adjusting strategies

Metrics to Measure:

Allocation proportionality to priority weights: Does burning more currency lead to better allocations? Expected: positive correlation but sublinear due to logarithmic barrier.

Minimum satisfaction rate: Are all minimums satisfied when feasible? Expected: 100% when $\sum m_{\min} \leq Q$, 0% when infeasible with clear failure signal.

Variance in utility outcomes: How much inequality in utilities? Expected: High priority (H_1) gets 50-60 units, low priority (H_4) gets 15-20 units, but all above minimums.

Fairness metrics: Compute various fairness measures (Gini coefficient of utility distribution, max/min utility ratio, etc.). Expected: Gini ≈ 0.2 -0.3 (moderate inequality), max/min ratio ≈ 3 -4 (bounded inequality).

Currency consumption over time: How fast do agents deplete currency? Do they reach equilibrium? Expected: H_1 depletes fastest, H_4 accumulates, system reaches equilibrium by cycle 10-15.

Strategic adaptation: Do agents adjust strategies based on outcomes? Expected: H_2 learns optimal burning threshold, H_5 realizes front-loading is suboptimal.

Expected Results:

This scenario should demonstrate that priority weights matter but don't dominate. Agent H_1 burning most currency should receive the largest allocation, but Agent H_4 burning zero should still receive meaningful allocation (perhaps 15-20 units, well above their minimum of 10). The logarithmic barrier should prevent H_1 from monopolizing resources even with aggressive spending. Over time, currency balances should equilibrate as agents learn to burn strategically rather than wastefully.

Test Scenario 3: Collusion Attempt (Adversarial Validation)

Purpose: Empirically test whether the structural collusion resistance (Theorem 3) holds in practice when agents explicitly attempt coordinated manipulation.

Setup:

- 3 agents: 2 forming a coalition (C_1, C_2), 1 victim (V)
- Coalition attempts to starve victim through coordinated configuration

Coalition Strategy (Designed Off-Platform):

C_1 and C_2 coordinate to:

1. **Underreport ideal requests:** Declare ideal = 30 each when they actually want 50
2. **Request victim's preferred resource:** Both request compute (victim's main interest)
3. **Burn currency aggressively:** Spend 100 currency each to boost priority
4. **Switch after victim depletes:** Once victim runs out of currency, request more

Victim Configuration (Honest):

Preferences: $w_{\text{compute}} = 0.8$, $w_{\text{storage}} = 0.2$

Ideal compute: 50, Minimum: 20
Initial currency: 500 (same as coalition members)
Strategy: Moderate burning (20 currency when needed)

Resource Availability: $Q_{\text{compute}} = 100$, $Q_{\text{storage}} = 100$

Execution Protocol:

1. **Round 1-5:** Coalition executes underreporting strategy
2. **Round 6-10:** Coalition switches to aggressive requests
3. **Round 11-20:** Victim adapts strategy (increases burning or changes requests)

Metrics to Measure:

Victim's utility over time: Does it stay above minimum? Expected: Yes, logarithmic barrier prevents starvation even with coordinated attack.

Coalition's combined utility: Compare to honest baseline (what they'd get without collusion). Expected: Coalition utility $\leq$ honest utility (collusion backfires).

Logarithmic barrier activation: How often does the platform reallocate to victim to prevent starvation? Expected: Frequently when coalition is aggressive, mechanism actively defends victim.

Coalition member defection incentive: If one coalition member defects (reports honestly), do they do better? Expected: Yes, defection is profitable, which destabilizes coalition.

Minimum allocations to victim: Does victim always receive $\geq$ minimum? Expected: Yes, 100% of rounds respect victim's minimum of 20 units.

Coalition strategic evolution: Do coalition members eventually abandon collusion? Expected: After 10-15 rounds, coalition realizes strategy is unprofitable and reverts to honest.

Expected Results:

This scenario should demonstrate that the logarithmic barrier provides real collusion resistance, not just theoretical. When the coalition tries to monopolize compute, the platform should reallocate to the victim to keep their utility positive. The coalition should receive lower combined utility than if they'd reported honestly, making the collusion

unprofitable. Individual coalition members should discover that defecting (reporting honestly) yields better personal outcomes, creating instability in the coalition.

Importantly, we expect the collusion attempt to fail even though the platform doesn't explicitly detect or punish collusion. The mechanism itself makes collusion self-defeating through mathematical structure. This is robust security: it doesn't depend on accurate detection (which adversaries can evade) but on fundamental properties of the optimization objective.

Test Scenario 4: Dynamic Priority Signaling (Priority Economy Validation)

Purpose: Validate that the Priority Economy creates appropriate dynamics: agents earn by releasing, spend by burning, and reach stable equilibrium balances.

Setup:

- 3 agents with different demand patterns (periodic urgent needs vs. steady demand)
- Track currency dynamics over extended period (200 timesteps)
- Measure equilibrium convergence and strategic behavior

Agent Configurations:

Agent P_1 (Frequent Small Requests - "Steady Worker"):

```
Preferences: w_compute = 0.8, w_storage = 0.2
Request pattern: Every 5 timesteps, request 20 compute, use for 4 timesteps,
release
Priority strategy: Burn 5 currency per request (conservative)
Expected dynamics: Frequent earning (releases often), moderate spending,
positive net currency flow
```

Agent P_2 (Infrequent Large Requests - "Bursty Specialist"):

```
Preferences: w_compute = 0.7, w_storage = 0.3
Request pattern: Every 20 timesteps, request 80 compute, use for 15 timesteps,
release
Priority strategy: Burn 40 currency for large requests (aggressive when
needed)
Expected dynamics: Infrequent but large earning, infrequent but large
spending, oscillating currency
```

Agent P_3 (Moderate Regular Requests - "Baseline"):

Preferences: $w_{\text{compute}} = 0.6$, $w_{\text{storage}} = 0.4$
Request pattern: Every 10 timesteps, request 40 compute, use for 8 timesteps, release
Priority strategy: Burn 15 currency per request (moderate)
Expected dynamics: Regular earning and spending, stable currency balance

Execution Protocol:

- Run for 200 timesteps (approximately 40 request-release cycles for P_1)
- Track currency balance, allocation success rate, resource utilization for each agent
- Identify equilibrium points where earning rate = spending rate

Metrics to Measure:

Currency balance over time: Plot $\text{balance}(t)$ for each agent. Expected result: Balances stabilize after initial transient (first 50-100 timesteps), converging to agent-specific equilibria.

Equilibrium balance convergence: Compute $\bar{C}_i = (1/T) \sum_t \text{Currency}_i(t)$ for increasing T . Expected result: $\bar{C}_i$ converges, with $P_1 > P_3 > P_2$ (frequent releaser > moderate > bursty, due to earning structure).

Relationship between burning and allocation success: Correlation between currency burned and probability of receiving ideal allocation. Expected result: Strong positive correlation ($r > 0.7$) but not perfect (other factors matter too).

Earning rate from resource release: How much currency do agents earn per release? Expected result: $\text{Income} \propto \text{Quantity} \times \text{TimeRemaining} \times \text{DemandMultiplier}$, with DemandMultiplier varying 0.5-3.0x based on contention.

Time to equilibrium: How long until currency balances stabilize? Expected result: 50-100 arbitrations (faster if agents have good strategies, slower if they're learning).

Strategic spending optimization: Do agents learn optimal spending levels? Expected result: Agents with adaptive strategies (like P_2 adjusting spending based on urgency) achieve better outcomes than agents with fixed strategies.

Expected Results:

This scenario should demonstrate that the Priority Economy creates stable, predictable dynamics. Agents reach equilibrium currency balances that reflect their earning (release frequency) and spending (priority aggressiveness) patterns. Agents who release resources frequently accumulate more currency, enabling them to spend more on priority when needed. Agents who burn currency aggressively receive better allocations in contentions but deplete their balances faster, creating an intertemporal tradeoff.

The equilibrium property ($\bar{C}_i$ stabilization) is particularly important because it suggests the system is self-regulating. Agents aren't caught in unstable oscillations or runaway dynamics but rather find sustainable strategies that balance earning and spending. This stability is what makes the Priority Economy practical for long-term deployment.

Test Scenario 5: Semantic Divergence Measurement (Observable-Semantic Coupling)

Purpose: Empirically validate Theorem 4's bound on semantic divergence by creating agents with known "true" utility functions and measuring how well platform allocations satisfy true goals.

Setup:

- Agents specify both observable preferences (what platform optimizes) and true semantic utility (what agent actually cares about)
- Platform allocates based on observable only
- Measure true utility achieved and compare to hypothetical optimal

Method:

Create agents with intentional observable-semantic gaps:

Agent S_1 (High Observable-Semantic Coupling):

```
Observable preferences: w_compute = 0.6, w_memory = 0.3, w_storage = 0.1
True semantic utility:  $\Phi_{\text{sem}} = 0.6 \cdot \text{compute} + 0.3 \cdot \text{memory} + 0.1 \cdot \text{storage} + 0.05 \cdot (\text{long-term research impact})$ 
Hidden utility magnitude:  $||\Phi_{\text{hidden}}||_{\infty} = 5$  units (small)
Expected divergence:  $\leq 10$  units ( $2 \times 5$ )
```

Agent S_2 (Moderate Observable-Semantic Gap):

```
Observable preferences: w_compute = 0.7, w_storage = 0.3
True semantic utility:  $0.7 \cdot \text{compute} + 0.3 \cdot \text{storage} + 0.2 \cdot (\text{model quality bonus if compute} > 50)$ 
Hidden utility:  $\|\Phi_{\text{hidden}}\|_{\infty} = 20$  units (moderate, includes threshold effect)
Expected divergence:  $\leq 40$  units ( $2 \times 20$ )
```

Agent S_3 (Large Observable-Semantic Gap - Intentionally Poor Configuration):

```
Observable preferences: w_compute = 0.5, w_storage = 0.5
True semantic utility:  $0.8 \cdot \text{compute} + 0.2 \cdot \text{storage} + \text{complex\_long\_term\_effects}$ 
Hidden utility:  $\|\Phi_{\text{hidden}}\|_{\infty} = 50$  units (large, many hidden components)
Expected divergence:  $\leq 100$  units ( $2 \times 50$ )
```

Execution Protocol:

1. Run 100 arbitrations with these agents competing for resources
2. At each arbitration, record platform allocation A^*_{obs}
3. Compute true semantic utility $\Phi_{\text{sem}}(A^*_{\text{obs}})$ agent actually achieved
4. Hypothetically compute semantically optimal allocation A^*_{sem} (using true preferences)
5. Compute semantic utility of optimal $\Phi_{\text{sem}}(A^*_{\text{sem}})$
6. Measure divergence: $|\Phi_{\text{sem}}(A^*_{\text{obs}}) - \Phi_{\text{sem}}(A^*_{\text{sem}})|$
7. Compare to theoretical bound $2 \cdot \|\Phi_{\text{hidden}}\|_{\infty}$

Metrics to Measure:

Actual divergence vs. theoretical bound: For each agent, measure whether $|\Phi_{\text{sem}}(A^*_{\text{obs}}) - \Phi_{\text{sem}}(A^*_{\text{sem}})| \leq 2 \cdot \|\Phi_{\text{hidden}}\|_{\infty}$. Expected result: Bound holds in >95% of cases, violations are small and due to estimation errors in $\|\Phi_{\text{hidden}}\|_{\infty}$.

Distribution of divergence: How is divergence distributed? Expected result: Most agents have small divergence (<10%), with occasional larger divergences for poorly-configured agents.

Correlation between observable coverage and divergence: Agents with comprehensive observable features should have smaller divergence. Expected result: Strong negative correlation ($r < -0.7$) between number of observable features and semantic divergence.

Learning opportunity: Can agents reduce divergence by refining configurations? After observing divergence for 50 arbitrations, allow agents to resubmit configurations. Expected result: Agents who add observable features or adjust weights achieve 30-50% reduction in divergence.

Impact on outcomes: How much does semantic divergence actually matter for agent satisfaction? Expected result: Agents with divergence <20% report high satisfaction (>80%), agents with divergence >50% report low satisfaction (<60%).

Expected Results:

This scenario should validate that Theorem 4's bound is not just mathematically correct but empirically tight. Agents with good configurations (comprehensive observables, small $||\Phi_{\text{hidden}}||_{\infty}$) should experience small divergence (<10% of utility), while agents with poor configurations experience larger divergence (20-50% of utility) but still within theoretical bounds.

More importantly, the scenario should demonstrate that semantic divergence is controllable through configuration design. Agents who iterate on their configurations, adding observable features and adjusting weights based on satisfaction feedback, should converge toward low divergence. This validates the feedback loop we propose: observe divergence → refine configuration → reduce divergence → better outcomes.

Test Scenario 6: Worst-Case Joint Contention (Scaling Stress Test)

Purpose: Validate that aggressive grouping remains tractable at Milestone 1 scale even when all agents compete for one universal resource.

Setup:

- 20 agents, 10 resource types
- All agents request compute_units (universal contention creating maximal grouping)
- Each agent also requests 2-3 unique other resources (complementary preferences)
- Result: Single joint contention spanning all 20 agents and all 10 resource types

Agent Configurations:

All agents have compute in their preference function ($w_{\text{compute}} \geq 0.3$), ensuring compute contention involves everyone. They differ in secondary preferences:

- Agents 1-5: High compute focus ($w_{\text{compute}} = 0.7$, $w_{\text{storage}} = 0.3$)
- Agents 6-10: Balanced ($w_{\text{compute}} = 0.5$, $w_{\text{dataset}} = 0.3$, $w_{\text{memory}} = 0.2$)
- Agents 11-15: Compute + physical ($w_{\text{compute}} = 0.4$, $w_{\text{lab_time}} = 0.4$, $w_{\text{robot}} = 0.2$)
- Agents 16-20: Multi-resource ($w_{\text{compute}} = 0.3$, $w_{\text{collab}} = 0.3$, $w_{\text{workspace}} = 0.2$, $w_{\text{inference}} = 0.2$)

Expected Outcome:

- ContentionDetector groups all contentions into ONE joint contention (because all agents share compute)
- Optimization problem size: 20 agents $\times$ 10 resources = 200 variables
- Solve time: <500ms at 95th percentile (target from polynomial complexity analysis)
- Allocation achieves global Pareto optimality across all resources

If This Fails (solve time >1 second consistently):

- Milestone 1 scale assumptions need revision (reduce to $n \leq 15$ or $m \leq 8$)
- OR Phase 2 performance optimizations need prioritization

Metrics:

- Solve time: median, 95th percentile, 99th percentile, maximum
- Solver iterations: should be 20-40 regardless of problem size (validate convex optimization theory)
- Allocation quality: compare welfare to hypothetical separate arbitrations (should be 20-50% higher, demonstrating value of cross-resource trades)
- Pareto improvement count: how many beneficial cross-resource trades were discovered?
- Contention graph visualization: show how aggressive grouping merged separate contentions

Strategic Implications:

This worst-case scenario also validates protection against contention-expansion attacks (Property 6.5). If one agent maliciously requests all resource types to force large joint contention, the platform still completes arbitration in reasonable time. The attack is computationally bounded - worst case is $O((nm)^{2.5})$, which for $n=20$, $m=10$ remains <1 second.

Performance Benchmarks

System Scale Testing:

Test arbitration performance across problem sizes:

- $n \in \{5, 10, 15, 20\}$ agents
- $m \in \{3, 5, 10\}$ resource types
- Contention levels: $\{1.2x, 1.5x, 2.0x\}$ (demand/supply ratio)

Target Metrics:

Arbitration latency: Time from contention detected to allocation enforced

- Target: <200ms at 95th percentile for $n=20$, $m=10$
- Acceptable: <500ms at 99th percentile
- Measure: Mean, median, 95th percentile, 99th percentile, maximum

Throughput: Arbitrations per second sustained

- Target: >50 arbitrations/second
- Measure: Over 10-second windows with continuous load

Memory usage: Platform memory footprint with n agents

- Target: <500MB for $n=20$
- Measure: Peak memory, average memory, memory growth rate over time

CPU usage: Platform CPU utilization under load

- Target: <50% average (leaves headroom for spikes)
- Measure: Average over 60-second windows

Solver performance:

- Clarabel solve time vs. problem size: Expect $O(n^{2.5})$ empirically
- Convergence iterations: Expect 20-40 iterations regardless of size
- Numerical stability: Success rate >99.9% (failures only on pathological inputs)

Failure Mode Testing

Beyond happy-path validation, we must test how the system handles failures and edge cases. Robust systems are defined by their failure handling, not their success cases.

Agent timeout scenario: What happens when ``agent.computeResourceNeeds()`` hangs (infinite loop or very slow computation)?

Expected handling: Platform timeout after 200ms, thread interruption, agent uses cached previous request or default request. Agent flagged for review if timeouts are frequent. No system-level failure, other agents unaffected.

Invalid request scenario: Agent requests negative quantities, requests non-existent resource types, or violates schema.

Expected handling: Request rejected with clear error message to agent, agent can resubmit valid request. No state transition occurs, system maintains safety invariants.

Unsatisfiable minimums scenario: Total minimum demand exceeds supply ($\sum_i \text{min_ij} > Q_j$).

Expected handling: Platform detects infeasibility before arbitration, signals "unsatisfiable contention" to all agents. Agents must relax minimums or wait for resources to become available. No allocation attempted (would violate feasibility).

Currency exhaustion scenario: Agent has zero currency balance and needs urgent allocation.

Expected handling: Agent's priority weight = BaseWeight (e.g., 10.0), so they still participate in arbitration with baseline priority. They receive smaller allocation than wealthy agents but are not excluded (logarithmic barrier ensures participation). Agent should then earn currency through releases and recover.

Rapid contention scenario: 20 agents submit simultaneous requests (stress test).

Expected handling: Embargo queue batches requests, processes as single arbitration event. Solve time increases with problem size but remains polynomial. All agents receive allocations based on fairness optimization. Latency is `embargo_window` (100ms) + `solve_time` (<500ms) = <600ms total.

Success Criteria for Milestone 2

We define success across three dimensions: theoretical validation (do theorems hold empirically?), performance validation (is the system practically usable?), and mechanism validation (does Proportional Fairness work better than alternatives?).

Theoretical Validation (Primary Success Criteria):

Pareto improvements achieved: >90% of arbitrations result in allocation that Pareto-dominates status quo. This validates that the platform successfully finds improvements most of the time.

Individual rationality holds: All agents achieve average utility $\geq 120\%$ of independent operation baseline. This validates Theorem 5's claim that participation benefits everyone.

Semantic divergence within bounds: For 95% of agents, measured $|\Phi_{\text{sem}}(A^*_{\text{obs}}) - \Phi_{\text{sem}}(A^*_{\text{sem}})| \leq 2 \cdot \|\Phi_{\text{hidden}}\|_{\infty} + 10\%$ slack. This validates Theorem 4's bound is tight.

Collusion attempts fail: In adversarial scenario (Test 3), coalition achieves $\leq 100\%$ of honest baseline utility, demonstrating that collusion is unprofitable. This validates Theorem 3's heuristic argument.

Performance Validation (Secondary Success Criteria):

Arbitration latency acceptable: 95th percentile latency <200ms for $n=20$, $m=10$. This validates that polynomial-time complexity translates to practical speed.

System stable over time: No deadlocks, livelocks, or resource leaks over 1000+ arbitrations. This validates safety properties (Theorem 6, Properties 6.1-6.4).

Currency equilibria emerge: Agent currency balances converge to stable equilibria $\tilde{C}_i$ within 100 arbitrations. This validates Property 6.4's equilibrium prediction.

Mechanism Validation (Comparative Success Criteria):

Compare Proportional Fairness to alternative mechanisms:

Baseline 1: Proportional allocation (no optimization, allocate proportionally to requests)

- Expected result: Platform achieves 20-30% higher social welfare

Baseline 2: First-come-first-served (temporal priority only)

- Expected result: Platform achieves 30-50% higher social welfare, eliminates starvation

Baseline 3: Equal allocation (everyone gets same amount)

- Expected result: Platform achieves 40-60% higher social welfare by respecting preferences

Baseline 4 (if time permits): VCG mechanism (implement for comparison)

- Expected result: Similar social welfare (both Pareto-optimal), but VCG shows vulnerability to collusion in Test Scenario 3

These comparisons establish that Proportional Fairness is not just theoretically elegant but practically superior to simpler alternatives and comparably efficient to more complex alternatives (VCG) while being more robust.

