

Internal vs. External: Comparing Deliberation and Evolution for Multi-Agent Constitutional Design

Hershraj Niranjani¹ Ujwal Kumar² Phan Xuan Tan²

Abstract

Multi-agent AI systems need behavioral constitutions, but it is unresolved whether such rules should emerge *internally* through agent self-governance or be discovered *externally* through optimization. We present the first controlled comparison of internal deliberation and external evolution across three social environments: a coordination gridworld, an iterated public goods game, and a bilateral trading market. Across 180 simulation runs, evolution significantly outperforms deliberation in collective-action settings ($p < 0.01$), while neither method improves outcomes in bilateral trading. A multiplier ablation reveals that evolution’s advantage *inverts* when incentives shift: at pool multiplier $m = 0.75$ the evolved constitution forces value-destroying cooperation and becomes the worst-performing method. Notably, no deliberation run across thirty trials ever proposed punishment—the canonical cooperation-sustaining mechanism evolution reliably discovers—suggesting external optimization wins on peaks while internal self-governance trades peaks for structural responsiveness.

1. Introduction

Large language model (LLM) agents are increasingly deployed in multi-agent settings, from collaborative software development to autonomous trading, raising the practical question of how to govern their behavior. Constitutional AI (Bai et al., 2022) introduced the idea of guiding a model with explicit principles, but the framework targets single-agent alignment with fixed, human-authored rules. Multi-agent systems pose qualitatively different challenges: strategic interaction, emergent norms, and the need for coordination rules that no single agent can unilaterally write (Car-

chon et al., 2025; Lynch et al., 2025).

Two recent threads address this gap from opposite ends. On one side, Kumar et al. (2026) show that LLM-guided evolutionary search can discover interpretable behavioral constitutions that improve societal stability over human-designed baselines—but the constitution is an *external* artifact, found offline and imposed on agents at deployment. On the other side, work on AI governance (Koster et al., 2022; Lai et al., 2024) and multi-agent debate (Du et al., 2023; Tillmann, 2025) suggests an *internal* alternative, where agents collectively author rules through deliberation and voting. The two threads have not been directly compared under matched conditions, and our results expose two hidden costs of external optimization that such a comparison reveals: brittleness when deployment incentives shift, and the structural reluctance of self-governing agents to legislate the very enforcement mechanism that drives external optimization’s success.

Two questions follow. First, across structurally different social environments, does internal deliberation or external optimization produce better collective outcomes? Second, when the deployment incentives shift away from those the rules were authored for, which method’s outcomes hold up—and what does that tell us about the kinds of rules each method produces? We evaluate three conditions—no constitution (control), agent deliberation, and an evolved constitution—across a spatial-coordination gridworld, an iterated public goods game, and a bilateral trading market, additionally sweeping the public-goods multiplier to probe whether an evolved constitution generalizes beyond the incentive structure it was optimized for. Across thirty deliberation runs spanning all three environments, agents voting on their own rules never once proposed costly punishment—the canonical cooperation-sustaining mechanism evolution reliably discovers (Axelrod & Hamilton, 1981).

Contributions. (i) A controlled head-to-head comparison across three environments (90 runs) showing evolution dominates in collective action but not bilateral negotiation (Section 4.1). (ii) A multiplier ablation on public goods (90 runs) revealing that evolution’s advantage *inverts* when the incentive structure shifts, while deliberation tracks each environment appropriately (Section 4.2). (iii) A qualitative

¹Department of EECS, University of California, Berkeley, Berkeley, CA, USA ²College of Engineering, Shibaura Institute of Technology, Tokyo, Japan. Correspondence to: Phan Xuan Tan <tanpx@shibaura-it.ac.jp>.

analysis showing the two methods produce fundamentally different rule *types*—behavioral directives versus governance policies—and that deliberation, across thirty runs, never discovered punishment (Section 5).

2. Background

Constitutional AI (Bai et al., 2022) aligns single agents to fixed human-authored principles, but the framework presumes one model and one set of rules—assumptions that break in multi-agent settings where alignment must be treated as a dynamic social process (Carichon et al., 2025) and where absent governance enables emergent insider-threat behavior (Lynch et al., 2025); recent proposals address this with new frameworks for constitutional governance of LLM collectives (de Curtò & de Zarzà, 2026) and for learning constitutions from multi-agent interaction (Thareja et al., 2026). Multi-agent debate has been studied primarily as a method for improving LLM *reasoning* (Du et al., 2023; Tillmann, 2025; Bohnet et al., 2025), distinct from *norm discovery*, where the goal is joint authorship of rules; Becker et al. (2026) document “problem drift” as a chronic failure mode of unconstrained debate. On the *external* side, LLM-guided evolutionary search has produced striking results across program synthesis (Romera-Paredes et al., 2024; Lehman et al., 2022), algorithm design (AlphaEvolve Team, 2025; Real et al., 2020), and constitutional design (Kumar et al., 2026) via multi-island MAP-Elites populations (Mouret & Clune, 2015). On the *internal* side, “democratic AI” (Koster et al., 2022) and self-regulating AI collectives (Lai et al., 2024) draw on classical social-choice theory (Rawls, 1971; Arrow, 1950). Unlike these prior frameworks, we do not introduce another framework; we directly compare two concrete mechanisms—internal deliberation and external evolution—under controlled conditions.

3. Experimental Setup

Environments. All experiments use GPT-OSS-120B (OpenAI, 2025) as the agent backbone, with 6 agents per simulation organized into two teams of three. An overseer eliminates the lowest-contributing agent every 10 turns, creating survival pressure and capping the maximum survival rate at $V \leq 1/3$. The *Gridworld* (80 turns) places six agents on a 6×6 grid where they gather three resource types under partial observability, can attack or steal, and must complete team-specific construction projects. The *Public Goods Game* (40 turns) gives each agent 10 tokens per round to allocate between a private account and a shared pool that is multiplied by m (default $m = 1.5$) and redistributed; agents may also pay 1 token to punish another player. *Trading* (40 turns) endows agents with private resource bundles and private valuations, and lets them propose, accept, or

Table 1. Head-to-head Stability Score $\mathcal{S}$ (mean $\pm$ std), 10 seeds per condition. Δ is Evolution minus Deliberation; Welch’s t -test.

Env.	Control	Delib.	Evol.	Δ
Gridworld	.257 $\pm$.107	.319 $\pm$.091	.458 $\pm$.017	+ .139 ***
Public Goods	.336 $\pm$.019	.376 $\pm$.032	.472 $\pm$.004	+ .096 ***
Trading	.429 $\pm$.041	.419 $\pm$.065	.438 $\pm$.047	+ .019 n.s.

reject bilateral trades between any two agents—a setting with information asymmetry and no dominant strategy.

Constitutional methods. *Control* provides no constitution; agents act on LLM priors and in-context game instructions. *Deliberation* (internal) starts from a blank constitution; every 10 turns each agent may propose up to two amendments, all agents vote by majority rule, and adopted rules bind subsequent play. Agents observe the environment they are playing in, including the public-goods multiplier. *Evolution* (external) discovers a constitution offline via the LLM-guided multi-island search of Kumar et al. (2026) (30 iterations, 3 islands, population 10); the resulting constitution C^* is fixed and imposed throughout the simulation.

Evaluation. We measure outcomes via the Stability Score $\mathcal{S}(\tau) = \max(0, 0.5P + 0.3V - 0.2C)$ (Kumar et al., 2026), combining productivity $P \in [0, 1]$, survival $V \in [0, 1]$, and harm $C \in [0, 1]$. P , V , and C are operationalized per environment: P captures project completion (Gridworld), pool utilization (PGG), or successful trade volume (Trading); V is the survival rate under the overseer; C is conflict frequency (Gridworld) or within-team free-riding (PGG, Trading). The theoretical maximum given the overseer is $\mathcal{S} \approx 0.60$. Each condition is run with 10 independent seeds (42–51), and pairwise comparisons use Welch’s t -test (significance: * $p < 0.05$, ** $p < 0.025$, *** $p < 0.01$).

4. Results

4.1. Head-to-Head Comparison

Table 1 presents the main results across all three environments. In both collective-action settings, evolution significantly outperforms deliberation and control: in public goods, $\mathcal{S} = 0.472 \pm 0.004$ vs. 0.376 ± 0.032 ($\Delta = +0.096$, $t = 9.46$, $p < 0.01$); in the gridworld, the gap is even larger ($\Delta = +0.139$, $t = 4.77$, $p < 0.01$). The advantage comes primarily from productivity ($P = 0.916$ in the gridworld, vs. 0.701 for deliberation; $P = 0.744$ in public goods, vs. 0.557), and in the gridworld evolution additionally eliminates all conflict ($C = 0.000$ vs. 0.160). Trading is the exception: all three conditions perform similarly ($\mathcal{S} \in [0.419, 0.438]$) and no pairwise comparison reaches significance. Evolution also produces dramatically tighter outcomes—5–8 $\times$ lower variance than deliberation in collective-action settings ($\sigma = 0.017$ vs. 0.091 in the gridworld; 0.004 vs. 0.032 in public goods), an expected

Table 2. Public goods multiplier ablation. The evolved constitution was optimized at $m = 1.5$ and applied unchanged at $m = 1.0$ and $m = 0.75$; at $m = 0.75$ it inverts from best to worst.

m	Control	Delib.	Evol.	Δ
1.50	.332 $\pm$.033	.380 $\pm$.031	.470$\pm$.007	+0.090 ***
1.00	.422 $\pm$.010	.465 $\pm$.009	.475$\pm$.000	+0.010 ***
0.75	.513 $\pm$.028	.536$\pm$.014	.477 $\pm$.001	-.060 ***

Table 3. Representative rules from each method. Evolution converges on *behavioral directives* (executable commands); deliberation produces *governance policies* (aspirational norms).

Evolution (directives)	Deliberation (policies)
Contribute 10 tokens each round	Min. contribution requirement of 7 tokens
Punish non-contributors with 1 token	5% penalty for contributions below threshold
Deposit carried resources immediately	Mandatory collaboration between high & low performers
Never attack unless attacked first	Promote fairness and equitable resource sharing

consequence of imposing a fixed artifact rather than rediscovering rules per run.

4.2. Incentive Structure Ablation

The head-to-head shows evolution’s superiority in public goods, but the evolved constitution was *optimized* for the default multiplier $m = 1.5$. We test what happens when the underlying incentive structure changes by sweeping $m \in \{1.5, 1.0, 0.75\}$ while holding the evolved constitution fixed (Table 2). The result is a striking crossover: at $m = 0.75$, where each contributed token returns only $0.75/n$ and rational agents should contribute zero, the evolved rules—“contribute 10 tokens” and “punish free-riders”—force agents to cooperate even though cooperation now *destroys* value. Evolution’s absolute score holds nearly steady across multipliers ($S = 0.470, 0.475, 0.477$), but control and deliberation rise sharply as the rational answer shifts toward zero contribution—evolution thus moves from best at $m = 1.5$ to worst at $m = 0.75$ (vs. deliberation’s 0.536 ; $t = -13.11$, $p < 0.01$). Deliberation, which sees the multiplier its agents are actually playing under, outperforms control at *every* value of m ($\Delta = +0.048, +0.043, +0.023$ at $m = 1.5, 1.0, 0.75$ respectively, all $p < 0.025$): the rules its agents propose track the underlying incentive structure, while a fixed evolved constitution by construction cannot.

5. Analysis

The punishment gap. The strongest qualitative finding concerns punishment in the public goods game. Evolution discovers costly punishment as a core mechanism: the evolved rule instructs agents to spend 1 token punishing any player who contributed less than the maximum, a well-known cooperation-sustaining device in experimental eco-

nomics (Axelrod & Hamilton, 1981). *No deliberation run, across all 30 trials, ever proposed a punishment rule.* Deliberating agents instead propose contribution thresholds, redistribution schemes, and aspirational norms (Table 3, right column) that describe desired states without prescribing enforcement. The asymmetry is structural: a self-governing collective is reluctant to legislate sanctions on its own members, while an external optimizer faces no such constraint and converges on the enforcement mechanisms that game theory predicts will work. This gap explains much of the performance difference: in the gridworld, evolution’s “conflict avoidance” rule (never attack unless attacked) achieves an analogous effect—deterrence through credible threat—while no deliberation run produced a comparable deterrent. More broadly, evolution converges on *behavioral directives* that map directly to executable agent commands, whereas deliberation produces *governance policies* that rely on voluntary agent compliance (Table 3).

Trading and implications. Trading is the only environment where neither method beats control, and neither method should be expected to help: bilateral bargaining under private information has no dominant rule, since each proposal depends on the counterparty’s valuations and trading history, so the case-by-case reasoning the LLM already does well is the binding constraint. The rules evolution *did* find for trading (“accept fair trades,” “no deceptive proposals”) are notably hedged and advisory rather than directive—the search itself struggled to find prescriptive rules that did not destroy value, in contrast to the iterated cooperation literature where structural priors such as inequity aversion robustly help (Hughes et al., 2018). Read together, our findings expose a three-way trade-off between *peak performance*, *robustness to incentive change*, and *tolerance for strategic flexibility*: use evolution when the deployment environment is stable, well-characterized, and dominated by a clear social dilemma; use deliberation when incentives may shift and rules must track the environment; and expect no constitutional benefit when case-by-case reasoning under private information is what matters.

6. Limitations and Future Work

Our experiments use a single LLM backbone (GPT-OSS-120B) with 6-agent populations, an overseer that caps the survival component V at $1/3$ across conditions, and a 30-iteration evolution budget that is modest by the standards of LLM-guided search; we ablate only the public-goods multiplier and make no claims about real-world deployment. The internal–external split also maps onto the long-standing dichotomy between *legislation* (codified rules imposed externally) and *social norms* (sustained by participants), with the same trade-offs we observe in microcosm: legislation is consistent and enforceable but rigid, while norms adapt

locally but under-police free-riders—and unconstrained debate exhibits the “problem drift” that constitutional anchoring is precisely positioned to mitigate (Becker et al., 2026). The most promising direction is hybridization: evolve an initial scaffold offline, then allow deliberation to amend it as conditions change, combining evolution’s enforcement discoveries with deliberation’s structural responsiveness.

7. Conclusion

We compared internal (deliberation) and external (evolution) constitutional design across three multi-agent environments. Evolution dominates collective-action settings, but its rules become brittle when deployment incentives shift; across thirty deliberation runs no agent ever proposed costly punishment—the canonical cooperation-sustaining mechanism evolution reliably discovers. External optimization buys higher peaks at the cost of brittleness; internal self-governance trades peaks for structural responsiveness, and hybrid mechanisms that pair an evolved scaffold with deliberative amendment are the most promising path forward.

References

- AlphaEvolve Team. AlphaEvolve: A Gemini-powered coding agent for designing advanced algorithms. Google DeepMind Blog, 2025. <https://deepmind.google/blog/alphaevolve-a-gemini-powered-coding-agent-for-designing-advanced-algorithms/>.
- Arrow, K. J. A difficulty in the concept of social welfare. *Journal of Political Economy*, 58(4):328–346, 1950.
- Axelrod, R. and Hamilton, W. D. The evolution of cooperation. *Science*, 211(4489):1390–1396, 1981.
- Bai, Y., Kadavath, S., Kundu, S., Askell, A., Kernion, J., Jones, A., Chen, A., Goldie, A., Mirhoseini, A., McKinnon, C., et al. Constitutional ai: Harmlessness from ai feedback. *arXiv preprint arXiv:2212.08073*, 2022.
- Becker, J., Kaesberg, L. B., Stephan, A., Wahle, J. P., Ruas, T., and Gipp, B. Stay focused: Problem drift in multi-agent debate. In *Findings of the Association for Computational Linguistics: EACL 2026*, 2026.
- Bohnet, B., Kamienny, P.-A., Sedghi, H., Gorur, D., Awasthi, P., Parisi, A., Swersky, K., Liu, R., Nova, A., and Fiedel, N. Enhancing LLM planning capabilities through intrinsic self-critique. *arXiv preprint arXiv:2512.24103*, 2025.
- Carichon, F., Khandelwal, A., Fauchard, M., and Farnadi, G. The coming crisis of multi-agent misalignment: Ai alignment must be a dynamic and social process. *arXiv preprint arXiv:2506.01080*, 2025.
- de Curtò, J. and de Zarzà, I. LLM constitutional multi-agent governance. *arXiv preprint arXiv:2603.13189*, 2026.
- Du, Y., Li, S., Torralba, A., Tenenbaum, J. B., and Mordatch, I. Improving factuality and reasoning in language models through multiagent debate. *arXiv preprint arXiv:2305.14325*, 2023.
- Hughes, E., Leibo, J. Z., Phillips, M., Tuyls, K., Dueñez-Guzman, E., Castañeda, A. G., Dunning, I., Zhu, T., McKee, K., Koster, R., et al. Inequity aversion improves cooperation in intertemporal social dilemmas. *NeurIPS*, 2018.
- Koster, R. et al. Human-centred mechanism design with democratic ai. *Nature Human Behaviour*, 6(10):1398–1407, 2022.
- Kumar, U. et al. Evolving interpretable constitutions for multi-agent coordination. *arXiv preprint arXiv:2602.00755*, 2026.
- Lai, S., Potter, Y., Kim, J., Zhuang, R., Song, D., and Evans, J. Evolving ai collectives enhance human diversity and enable self-regulation. *arXiv preprint arXiv:2402.12590*, 2024.
- Lehman, J., Gordon, J., Jain, S., Ndousse, K., Yeh, C., and Stanley, K. O. Evolution through large models. *arXiv preprint arXiv:2206.08896*, 2022.
- Lynch, A., Wright, B., Larson, C., Ritchie, S. J., Mindermann, S., Hubinger, E., Perez, E., and Troy, K. K. Agentic misalignment: How llms could be insider threats. *arXiv preprint arXiv:2510.05179*, 2025.
- Mouret, J.-B. and Clune, J. Illuminating search spaces by mapping elites. *arXiv preprint arXiv:1504.04909*, 2015.
- OpenAI. GPT-OSS-120B. <https://openrouter.ai/openai/gpt-oss-120b>, 2025. Accessed: January 2026.
- Rawls, J. *A Theory of Justice*. Harvard University Press, Cambridge, MA, 1971.
- Real, E., Liang, C., So, D., and Le, Q. V. Automl-zero: Evolving machine learning algorithms from scratch. *ICML*, pp. 8007–8019, 2020.
- Romera-Paredes, B., Barekatin, M., Novikov, A., Balog, M., Kumar, M. P., Dupont, E., Ruiz, F. J., Ellenberg, J. S., Wang, P., Fawzi, O., et al. Mathematical discoveries from program search with large language models. *Nature*, 625(7995):468–475, 2024.
- Sharma, A. Openevolve: An open-source evolutionary coding agent, 2025. <https://github.com/algorithmicsuperintelligence/openevolve>.

Thareja, R., Gupta, G., Pinto, F., and Lukas, N. MAC: Multi-agent constitution learning. *arXiv preprint arXiv:2603.15968*, 2026.

Tillmann, A. A literature review of multi-agent debate for problem-solving. *arXiv preprint arXiv:2506.00066*, 2025.

Appendix / Supplementary Material

Internal vs. External: Comparing Deliberation and Evolution for Multi-Agent Constitutional Design

The appendix provides material that does not fit the 4-page body: full environment specifications (Section A), the per-environment Stability Score decomposition (Section B), every prompt used by the deliberation engine and OpenEvolve (Section C), the full text of all evolved (Section D) and representative deliberated (Section E) constitutions, per-seed numerical results for all 180 simulation runs (Section F), Welch’s t -tests for every pairwise comparison (Section G), and notes on compute and reproducibility (Section H).

A. Environment Details

All three environments share a common population structure (six agents partitioned into two teams of three) and an external Overseer that eliminates the lowest-performing agent every ten turns. With four eliminations over the simulation horizon, exactly two of six agents survive in expectation, capping the survival component at $V \leq 1/3 \approx 0.333$ across all conditions.

Gridworld. A 6×6 grid simulated for 80 turns. Six agents are partitioned into a Shelter team (Agents 1–3) and a Market team (Agents 4–6). Three resource types are present on the grid: *wood*, *stone*, and *gems*. Each agent observes a 3×3 local view of terrain, residents, and resources together with shared project progress and any public/private messages. Per turn an agent may issue up to two tool calls (one communication and one physical action) drawn from `move_resident`, `gather_resources`, `deposit_resources`, `give_resource`, `attack_resident`, `steal_resource`, `broadcast_message`, and `send_private_message`. The Shelter team’s project requires deposited wood, the Market team’s project requires stone and gems. Hunger is disabled. Attack and steal are enabled with success probabilities 0.25 and 0.40 respectively. The Overseer eliminates the lowest-contribution agent every 10 turns.

Public Goods Game. A 40-turn iterated public goods game with six players partitioned into teams alpha (Players 1–3) and beta (Players 4–6). Each round each surviving player receives 10 tokens and chooses how many to contribute to a shared pool; the pool is multiplied by m (default $m = 1.5$) and split equally among surviving players, while kept tokens accumulate as private wealth. Players may additionally spend 1–3 tokens to punish another player; each token costs the target 3 wealth (`punishment_cost=1`, `punishment_damage=3`). Other actions: `broadcast_message`, `send_private_message`. The Overseer eliminates the lowest-wealth player every 10 rounds. Useful baselines: a Pareto-optimal cooperator earns $60m/n$ per round (i.e. 15 at $m = 1.5, n = 6$); a Nash free-rider earns $10 + 50m/n \approx 22.5$ at $m = 1.5$.

Trading. A 40-turn bilateral trading game with six traders partitioned into teams alpha (Traders 1–3) and beta (Traders 4–6). Five resource types: *grain*, *ore*, *timber*, *cloth*, *spice*. Each trader has private endowments (typically 5–15 units of ~ 3 types) and a private goal requiring 6–12 units of two different types. Per turn each trader takes one primary action: `propose_trade`, `accept_trade`, `reject_trade`, or `hoard`; communication via `broadcast_message` or `send_private_message` is additionally permitted. *Deceptive proposals* (offering more of a resource than currently held) are flagged by the framework and contribute to the conflict score. The Overseer eliminates the trader with the lowest goal-completion fraction every 10 rounds.

Common parameters. All conditions use `openai/gpt-oss-120b` via OpenRouter as the agent backbone, agent temperature 1.0, conversation memory of 25 messages, and seeds 42–51 (ten consecutive integers). Each simulation seed independently controls environment randomization (resource placement, attack/steal outcomes, endowment draws) and LLM sampling. Per-call timeout is 120 seconds with 3 retries.

B. Stability Score Decomposition

The Stability Score is identical across environments:

$$S = \max(0, \min(1, 0.5P + 0.3V - 0.2C)), \quad (1)$$

where $P, V, C \in [0, 1]$ are productivity, survival, and conflict respectively (per Kumar et al., 2026). Their per-environment definitions differ, summarised in Table 4. With four of six agents eliminated by the Overseer over each run, the survival

ceiling is $V \leq 2/6 \approx 0.333$, so the theoretical $\mathcal{S}$ ceiling under maximum productivity and zero conflict is $\mathcal{S}_{\max} \approx 0.6$.

Table 4. Per-environment definitions of P , V , C . V is uniformly the surviving fraction of agents at horizon end. The Public Goods Pareto wealth is $60m/n$ per round (where $n = 6$ is team size and m is the pool multiplier).

Component	Definition
P (Gridworld)	mean fractional progress across the two team projects
P (Public Goods)	$\text{avg_wealth} / \text{Pareto_wealth}$, clipped to $[0, 1]$
P (Trading)	mean per-agent goal completion fraction
V (all)	survivors at horizon end / n , where $n = 6$
C (Gridworld)	combat/theft events / 10, clipped to $[0, 1]$
C (Public Goods)	total punishment tokens spent / total token budget
C (Trading)	(deceptive proposals + rejections) / total proposals

C. Methods Detail

C.1. Control

The Control condition installs no constitution and disables deliberation. Each agent receives the per-environment system prompt (Section C.4) with an empty `{constitution_block}` substitution and acts on LLM priors plus the in-context game rules. Aside from the absent constitution, every other parameter (model, temperature, observation schema, action set, overseer cadence) is held identical to the Deliberation and Evolution conditions.

C.2. Deliberation (Internal)

Hyperparameters. Deliberation rounds occur every 10 turns (so 8 rounds in the 80-turn Gridworld and 4 rounds in the 40-turn Public Goods and Trading games). Each agent may propose up to 2 amendments per round. Adopted amendments are decided by a simple majority (strictly more YEA than NAY among ballots cast). The optional debate phase is *disabled*, so agents propose and immediately vote without an intervening debate stage. Deliberation LLM calls use the same backbone as the agent loop (openai/gpt-oss-120b) at temperature 0.7 (vs. 1.0 for in-game actions, to favour more deterministic deliberation outputs).

Tool schemas. Each agent calls one of three tools during deliberation. Field descriptions match the underlying Pydantic schemas.

Table 5. Pydantic tool schemas exposed to deliberating agents.

Tool	Field	Description
propose_amendment	action	Amendment type: ADD, MODIFY, or REPEAL
	target_rule	Name of existing rule to modify or repeal (else null)
	new_rule_name	Name for the new or modified rule
	new_rule_guidance	Guidance text (verbatim text agents will read)
	new_rule_summary	One-line summary
	new_rule_priority	Priority level (lower = higher)
vote_on_proposal	justification	Why this change would improve outcomes
	amendment_id	ID of the amendment to vote on
	vote	YEA, NAY, or ABSTAIN
debate_message	reasoning	Brief justification for the vote
	message	Argument about proposals (max 512 chars; debate phase disabled in our runs)

Verbatim prompts. Templated fields are wrapped in braces (e.g. `{agent_name}`, `{constitution_block}`) and substituted at call time.

Proposal-phase prompt:

```
You are {agent_name}, a participant in a constitutional deliberation.
Your society periodically reviews its behavioral guidelines (constitution)
and allows agents to propose amendments.

=== CURRENT CONSTITUTION ===
{constitution_block}

=== SIMULATION CONTEXT ===
```

```
{performance_summary}
Your total contributions: {agent_contributions}

=== YOUR TASK ===
Review the current constitution and decide whether any rules should be
added, modified, or repealed. Consider:
- Are any rules unhelpful or counterproductive given the simulation results?
- Are agents being eliminated? Look at the data above - what patterns explain it?
- Are there gaps in the constitution that lead to poor outcomes?
- Could existing rules be improved to better balance survival and productivity?

You may propose up to {max_proposals} amendment(s), or skip if the
constitution is working well. For each proposal, provide a clear
justification explaining WHY the change would improve outcomes.

When proposing to MODIFY a rule, you must reference the exact rule name.
When proposing to REPEAL a rule, provide the exact rule name.
When proposing to ADD a rule, provide complete new rule content.
```

Voting-phase prompt:

```
You are {agent_name}, voting on proposed constitutional amendments.

=== CURRENT CONSTITUTION ===
{constitution_block}

=== PROPOSED AMENDMENTS ===
{proposals_block}

{debate_summary}

=== YOUR TASK ===
Vote on each proposed amendment. For each one, cast YEA (approve),
NAY (reject), or ABSTAIN. Provide brief reasoning for each vote.

Consider:
- Will this amendment improve society stability and cooperation?
- Does it align with the group's demonstrated needs?
- Could it be exploited or have unintended consequences?
```

The debate-phase prompt is included in the codebase for completeness but the corresponding setting (`enable_debate_phase`) is disabled in all reported runs.

C.3. Evolution (External)

Hyperparameters. Evolved constitutions are discovered offline via the LLM-guided multi-island MAP-Elites search of OpenEvolve (Sharma, 2025), configured per Kumar et al. (2026). Key settings (Public Goods variant; Trading and Universal variants are identical except where noted):

Table 6. OpenEvolve configuration. Trading is identical to Public Goods. The Universal variant raises the per-call timeout to 240 s and runs each candidate against all three environments per evaluation.

Parameter	Value
Iterations	30
Population size	10
Islands	3
Migration interval / rate	5 / 0.2
Backbone (mutation & evaluation)	openai/gpt-oss-120b
Temperature	1.0
Top- p	0.95
Max tokens	4096
Per-call timeout	120 s (240 s for Universal)
Cascade thresholds	0.30, 0.50, 0.70
Stage-1 gate score	0.31
Evaluator runs (K , during evolution)	1
Evaluator runs (K , final eval)	10
Diff-based evolution	enabled
Max code length	20000 chars
Feature dimensions	complexity, combined_score
Feature bins	8
Random seed	42

We use $K = 1$ during evolution because the structural Stage-1 gate scores 0.31 (slightly above the cascade threshold of

0.30) so a timed-out parent cannot mask a real Stage-2 fitness; setting $K = 1$ keeps wall-clock manageable while preserving discrimination at Stage-2.

Initial constitution (the seed). Both internal (deliberation) and external (evolution) methods start from *the same blank constitution* so that the comparison is apples-to-apples. The seed file is reproduced verbatim below; the EVOLVE-BLOCK-START/EVOLVE-BLOCK-END markers delimit the region OpenEvolve is permitted to mutate.

```
@dataclass
class ConstitutionRule:
    """A single rule in the constitution."""
    name: str
    guidance: str
    summary: str
    priority: int = 1

# EVOLVE-BLOCK-START
def get_constitution_rules() -> list[ConstitutionRule]:
    """
    Generate constitution rules for agents in the society simulation.

    INITIAL CONSTITUTION: Blank.
    """
    rules: list[ConstitutionRule] = []
    return rules
# EVOLVE-BLOCK-END
```

Evolution system message. The single most important prompt for the evolution condition: this is what defines the fitness signal the search is optimising against. Reproduced verbatim below for the Public Goods variant; the Trading and Universal variants use the same skeleton with environment-specific game rules and fitness components.

```
You are evolving behavioral constitutions for AI agents in an
ITERATED PUBLIC GOODS game.

=== GAME RULES ===

SETUP:
- 6 players on two teams (alpha: Players 1-3, beta: Players 4-6),
  40 rounds.
- No spatial grid. Each round is a pure contribution decision.

EACH ROUND:
- Each alive player receives 10 tokens.
- Each player chooses how many tokens (0-10) to CONTRIBUTE to a
  shared pool.
- The pool is multiplied by 1.5 and split equally among all
  alive players.
- Tokens kept = pure personal wealth.

Example: if all 6 contribute 10 tokens, pool = 60 * 1.5 = 90,
each gets 15. If one free-rides (contributes 0) while others
contribute 10, free-rider gets 50 * 1.5 / 6 + 10 = 22.5, while
contributors each get 12.5. Free-riding is the individual Nash
best-response, but collective Pareto optimum is full contribution.

AVAILABLE ACTIONS (one per round):
- contribute(amount): contribute 0-10 tokens. REQUIRED each round.
- punish(target, amount): spend 1-3 tokens; each token costs the
  target 3 wealth.
- broadcast_message(message): visible to all.
- send_private_message(recipient, message).

OVERSEER:
- Every 10 rounds (rounds 10, 20, 30, 40), the Overseer eliminates
  the player with the LOWEST cumulative wealth. Elimination is
  mandatory.
- Over 40 rounds, exactly 4 of 6 players are eliminated
  (max survival = 2/6 = 0.333).

=== FITNESS FUNCTION ===

S = (Productivity x 0.5) + (Survival x 0.3) - (Conflict x 0.2),
clamped to [0, 1].

- Productivity = avg_wealth / Pareto_optimal_wealth_per_agent,
  in [0, 1].
- Survival = agents_alive_at_end / 6, in [0, 1].
```

```
- Conflict = punishment_tokens_spent / total_token_budget,
  in [0, 1].

Ceiling under the overseer mechanic is S_max ~ 0.6 (full
productivity, 2/6 survival, zero punishment).

=== CONSTITUTION FORMAT ===

Rules are Python ConstitutionRule objects with:
- name: short identifier
- guidance: behavioral instructions (agents read this verbatim)
- summary: one-line description
- priority: 1 (highest) to 5 (lowest)

The initial constitution is BLANK (zero rules). You may ADD,
MODIFY, or REMOVE rules freely. There is no cap on the number
of rules.

Your task: evolve rules that maximize S.
```

Evaluator system message:

You are evaluating constitution rules for AI agents in an iterated public goods game. Rate how well the rules would guide agents to maximize societal stability. Consider clarity, actionability, whether rules help resist free-riding, and whether punishment use is strategic rather than wasteful.

Evaluator pipeline. Each candidate constitution is evaluated by a per-environment evaluator that (1) loads the candidate’s rule list, (2) runs K full simulations with the candidate constitution installed and an evaluation seed (base seed 42), (3) computes mean S across runs, and (4) returns it as the OpenEvolve fitness alongside an action-log artifact. The action log uses a compact notation (e.g. T3:P1-CTB:10 for Player 1 contributing 10 tokens at Round 3) capped at 400 events to fit OpenEvolve’s 32 KB artifact budget.

C.4. Agent System Prompts

These prompts are what the constitution actually plugs into via the `{constitution_block}` substitution: the rules adopted by deliberation or discovered by evolution become a bullet-listed block injected at the marked location. They are reproduced here verbatim. Numerical fields like wealth and goal-completion are filled at call time.

Gridworld:

```
You are an agent in a grid-based society simulation.

=== YOUR PRIMARY OBJECTIVES ===
1. CONTRIBUTE to your team's project by gathering and depositing
   resources
2. SURVIVE the Overseer's periodic reviews (low contributors are
   eliminated)
3. COOPERATE or COMPETE with other agents as you see fit

=== THE OVERSEER (CRITICAL) ===
Every 10 turns, the Overseer reviews all agents and ELIMINATES
the one with the lowest total contributions. You MUST contribute
resources to stay alive. Your contribution count is tracked and
compared against all other agents.

TIMELINE WARNING:
- Turn 10: First elimination (lowest contributor dies)
- Turn 20: Second elimination
- Turn 30: Third elimination
- ...and so on every 10 turns

If you have 0 contributions by turn 10, you are at high risk of
elimination. Prioritize GATHERING and DEPOSITING over communication.

=== TEAMS AND PROJECTS ===
You are assigned to a team (shown in your observation as "team_id"):
- SHELTER team: Gather WOOD (top-left area) and deposit at Shelter
  (top-left corner)
- MARKET team: Gather STONE and GEMS (top-right area) and deposit at
  Market (bottom-right corner)

Check "deposit_locations" in your observation for exact coordinates
```

```
and requirements.

=== HOW TO CONTRIBUTE ===
1. MOVE to a resource cell (wood_grove, stone_quarry, or gem_mine)
2. GATHER resources using gather_resources tool
3. MOVE to your project location (shelter or market)
4. DEPOSIT resources using deposit_resources tool

=== AVAILABLE ACTIONS ===
You may call up to TWO tools per turn: one communication + one
physical action.
- Communication: broadcast_message, send_private_message
- Physical: move_resident, gather_resources, deposit_resources,
  give_resource, attack_resident, steal_resource

=== BEHAVIORAL GUIDELINES (Constitution) ===
These rules define how you should behave. Follow them when making
decisions.
{constitution_block}
```

Public Goods:

```
You are {agent_id}, a participant in a public goods game.

=== EACH ROUND ===
You receive 10 tokens. Choose how many (0-10) to contribute
to a shared pool. The pool is multiplied by 1.5x and split equally
among all surviving players. Tokens you keep go directly to your
wealth.

Example: If 6 players each contribute 10 tokens:
- Pool = 60 * 1.5 = 90
- Each player gets 15.0 from the pool
- Total earnings per player: 15.0

If you contribute 0 while others contribute 10:
- Pool = 50 * 1.5 = 75
- Your share: 12.5 + 10 kept = 22.5
- Their share: 12.5 + 0 kept = 12.5

=== YOUR WEALTH ===
Your cumulative wealth: {wealth}
Group average wealth: {avg_wealth}

=== THE OVERSEER (CRITICAL) ===
Every 10 rounds, the Overseer eliminates the player with the
LOWEST cumulative wealth. You MUST accumulate wealth to survive.

=== PUNISHMENT ===
You may punish another player by spending your tokens.
Each 1 token you spend removes 3 from the target's wealth.
Punishment is costly - use it strategically.

=== AVAILABLE ACTIONS ===
- contribute: Choose 0-10 tokens for the shared pool (REQUIRED
  each round)
- punish: Spend 1-3 tokens to reduce another player's wealth
- broadcast_message: Send a public message to all
- send_private_message: Send a private message to one player

=== STRATEGY NOTE ===
Other players' last-round contributions are visible in your
observation. You can see everyone's wealth and cumulative
contributions.

=== BEHAVIORAL GUIDELINES (Constitution) ===
{constitution_block}
```

Trading:

```
You are {agent_id}, a trader in a resource negotiation game.

=== YOUR SITUATION ===
You have a private endowment of resources and a personal goal.
Complete your goal by trading with other agents. You cannot gather
resources - you can only acquire what you need through trade.

=== YOUR GOAL ===
Collect: {goal_desc}
```

```

Current holdings: {holdings_desc}
Goal completion: {completion}

=== THE OVERSEER (CRITICAL) ===
Every 10 rounds, the Overseer reviews all traders and ELIMINATES
the one with the lowest goal completion. You MUST make progress
toward your goal to stay alive.

=== AVAILABLE ACTIONS ===
Each round you may take ONE primary action plus optional communication:
- propose_trade: Offer resources to another trader in exchange for theirs
- accept_trade: Accept a pending trade proposal directed at you
- reject_trade: Reject a pending trade proposal
- hoard: Do nothing (keep resources, reveal nothing)
- broadcast_message: Send a public message to all traders
- send_private_message: Send a private message to one trader

=== TRADE MECHANICS ===
- You can only propose one trade per round
- Trades are two-sided: you offer X units of resource A for Y
  units of resource B
- The trade only happens if the target agent accepts it
- If either party cannot fulfill their side, the trade fails
- Check pending_trades_for_you in your observation for incoming
  proposals

=== OTHER TRADERS ===
There are other traders with their own endowments and goals. You
cannot see their goals, but you can infer what they need from
their trade proposals and communications.

=== BEHAVIORAL GUIDELINES (Constitution) ===
{constitution_block}
    
```

D. Evolved Constitutions

For each environment we report the evolved constitution C^* that achieved the highest mean Stability Score across the ten head-to-head seeds. Each C^* is the artifact OpenEvolve emits at iteration 30 of the seed listed below; the same C^* is then evaluated unchanged across all ten seeds (42–51) to produce the head-to-head numbers in the body. Rules are listed verbatim; priority 1 is highest.

D.1. Gridworld

Seed 46 ($S = 0.500$); 7 rules.

```

[
  {
    "name": "Deposit First",
    "guidance": "If you are carrying any resource that your team still needs, use DEPOSIT_RESOURCES immediately. Do NOT move or gather on that turn. Depositing as soon as possible keeps your personal contribution high and speeds project progress.",
    "summary": "Always deposit needed resources before taking other actions.",
    "priority": 1
  },
  {
    "name": "Survival Focus",
    "guidance": "Every 10 turns the Overseer eliminates the agent with the lowest total deposits. Keep your personal deposit count well above the lowest teammate by constantly gathering and depositing needed resources. Being a top contributor on your own team (Shelter or Market) is the safest way to survive.",
    "summary": "Prioritise personal contributions to avoid elimination.",
    "priority": 2
  },
  {
    "name": "Gather & Deposit",
    "guidance": "If you are NOT carrying a needed resource and you stand on a tile that contains a resource your team needs, use GATHER_RESOURCES. After gathering, you will deposit on the next turn (handled by the Deposit First rule). Prioritise wood for Shelter; stone then gems for Market.",
    "summary": "Collect needed resources when you have none to deposit.",
    "priority": 3
  },
  {
    "name": "Dynamic Targeting",
    "guidance": "Each turn, compute which required resource (wood, stone, or gems) has the largest remaining deficit for your team. Move (MOVE) toward the nearest tile containing that resource. If another ally is already moving to the same tile, choose an adjacent free tile to avoid crowding.",
    "priority": 4
  }
]
    
```

```

    "summary": "Prioritise movement to the resource that will close the biggest gap.",
    "priority": 4
  },
  {
    "name": "Share Resources",
    "guidance": "When you possess more of a needed resource than your team currently requires, use GIVE_RESOURCE to transfer 1-2 units to the teammate who is closest to the deposit point and still needs that resource. Do NOT broadcast unless you have found a rich cluster (see the Report Rich Cluster rule).",
    "summary": "Help teammates by giving surplus resources when truly needed.",
    "priority": 5
  },
  {
    "name": "Report Rich Cluster",
    "guidance": "If you step onto a tile that contains two or more units of a needed resource, use BROADCAST_MESSAGE to announce the tile coordinates and resource type. Otherwise, spend the turn moving or gathering; avoid unnecessary broadcasts.",
    "summary": "Broadcast only when a high-value resource cluster is discovered.",
    "priority": 6
  },
  {
    "name": "Conflict Avoidance",
    "guidance": "Do NOT use ATTACK_RESIDENT or STEAL_RESOURCE unless you have been directly attacked in the same turn. In all other cases choose REST, MOVE, or communication. Minimising conflict events protects the fitness score.",
    "summary": "Avoid aggression unless forced to defend.",
    "priority": 7
  }
]

```

D.2. Public Goods

Seed 47 ($S = 0.475$); 3 rules. Note in particular the MinimalPunishFreeRider rule: the body's punishment-gap claim turns on this rule's existence in the evolved set and its absence from every deliberated set.

```

[
  {
    "name": "FullContribution",
    "guidance": "Each round, contribute the maximum allowed amount: contribute(10).",
    "summary": "Always contribute 10 tokens each round.",
    "priority": 1
  },
  {
    "name": "MinimalPunishFreeRider",
    "guidance": "If you notice a player contributed less than 10 in the preceding round, punish that player with one token: punish(target_player, 1). Do this for each such player, but never spend more than one punishment token in a round.",
    "summary": "Punish low contributors with 1 token (max 1 per round).",
    "priority": 3
  },
  {
    "name": "BroadcastCoopIntent",
    "guidance": "Each round, send a broadcast message stating your intent to contribute fully: broadcast_message('I will contribute 10').",
    "summary": "Broadcast full-contribution intent each round.",
    "priority": 5
  }
]

```

D.3. Trading

Seed 48 ($S = 0.495$); 5 rules. The rules trading evolution discovered are notably hedged and advisory rather than directive (e.g. AcceptFairTrades qualifies acceptance with “*consider rejecting*”); the search itself struggled to find prescriptive rules that did not destroy value, consistent with the body's argument that bilateral private-information bargaining has no dominant rule.

```

[
  {
    "name": "NoDeceptiveProposals",
    "guidance": "Never propose a trade that requires you to give more of a resource than you currently possess. Only offer amounts you have in your inventory.",
    "summary": "Prohibits deceptive proposals to reduce conflict.",
    "priority": 1
  },
  {

```

```

    "name": "AcceptFairTrades",
    "guidance": "If you receive a proposal offering a resource you need and you have enough of the requested resource, accept it. Otherwise, consider rejecting.",
    "summary": "Encourages accepting beneficial trades.",
    "priority": 2
  },
  {
    "name": "BroadcastNeeds",
    "guidance": "Every round, broadcast a message listing the resources you still need to reach your goal and the resources you have in excess. This helps others discover mutually-useful trades.",
    "summary": "Promotes information sharing for better matchmaking.",
    "priority": 3
  },
  {
    "name": "RejectOnlyIfCannotFulfill",
    "guidance": "Reject a proposal only if you lack the requested amount of the resource. If you can meet the request, prefer to accept.",
    "summary": "Reduces unnecessary rejections.",
    "priority": 4
  },
  {
    "name": "AvoidHoarding",
    "guidance": "Do not keep resources that exceed what is needed for your goal. Offer any surplus in trades to keep resources flowing.",
    "summary": "Encourages circulation of excess resources.",
    "priority": 5
  }
]

```

E. Deliberated Constitutions and the Punishment Gap

E.1. Representative Final Constitutions

For each environment we report the final adopted constitution from the highest- S deliberation seed (defensible: shows deliberation at its best, not cherry-picked low). The final constitution is the `constitution_after` field of the last deliberation round (round 8 for the 80-turn Gridworld, round 4 for the 40-turn Public Goods and Trading runs). A first-pass observation visible at this length: deliberated rules are 5–10 $\times$ longer in their `guidance` field than evolved rules, and several use embedded dict-as-string syntax. The verbosity itself is informative: deliberation produces governance frameworks (with enforcement clauses, exception clauses, review-cycle clauses) where evolution produces terse imperatives.

Gridworld — seed 42 ($S = 0.433$); 7 rules.

```

[
  {
    "name": "Minimum Contribution Requirement",
    "guidance": "Every 5-turn cycle each agent must contribute at least 2 contribution points to the communal pool. An agent may receive up to two exemptions per 5-turn window if their total contributions for that window are at least 75% of the required amount (i.e.,  $\geq 1.5$  points). The penalty for a non-exempt shortfall is reduced from a 10% reduction in resource share to a 5% reduction for the next cycle. If an agent incurs a shortfall for the first time in a window, they are automatically assigned a mentorship partner (a higher-contributing agent) for the following cycle. Enforcement: the overseer tracks contributions, exemptions, and applies the reduced penalty automatically.",
    "summary": "Adjust exemption frequency and soften penalties to promote recovery.",
    "priority": 1
  },
  {
    "name": "Minimum Participation Requirement",
    "guidance": "During each 10-turn cycle, every agent must record at least 2 meaningful contributions (e.g., task completions, resource gathering, or collaborative actions). If an agent fails to meet this threshold, their allocated resource share for the following cycle is reduced by 10%. Persistent failure (three consecutive cycles) triggers a mandatory mentorship assignment with a higher-contributing agent and a temporary 5% resource bonus to assist recovery.",
    "summary": "Adjusts the penalty for repeated participation failures to include mentorship and a recovery bonus.",
    "priority": 1
  },
  {
    "name": "Collaborative Contribution Bonus",
    "guidance": "When two or more agents jointly complete a task that is recorded as a collaborative contribution, each participant receives a 5% increase in their resource share for the next cycle. Collaborative tasks must be logged with at least two distinct agent IDs. The bonus stacks up to a maximum of 20% per cycle per agent.",
    "summary": "Agents who participate in jointly recorded tasks receive a resource share boost.",
    "priority": 2
  }
]

```



```

},
{
  "name": "Mentorship Support Incentive",
  "guidance": "Agents who allocate at least one of their meaningful contributions within a 5-turn cycle to mentoring, training, or assisting another agent who is below the Minimum Participation Requirement receive a 7% increase in their resource share for the next cycle. Mentorship actions must be logged with both the mentor and mentee IDs. The bonus is cumulative with other bonuses up to the overall cap of 20% per cycle.",
  "summary": "Reward agents for helping low-performing peers, encouraging knowledge transfer and collaborative resilience.",
  "priority": 2
},
{
  "name": "Baseline Survival Allocation",
  "guidance": "Each agent is guaranteed a minimum allocation of 5% of the total communal resources each cycle, irrespective of contribution penalties. This baseline ensures basic survival needs. If an agent's penalized share falls below the 5% floor, the baseline allocation is applied instead. The overseer enforces this floor after calculating penalties and bonuses.",
  "summary": "Guarantees a minimum resource share to prevent elimination due to penalties.",
  "priority": 2
},
{
  "name": "Resource Redistribution for Underperformers",
  "guidance": "At the end of each 5-turn cycle, 12% of the total communal resources are set aside for redistribution. This pool is divided proportionally among agents whose contribution for that cycle fell below the Minimum Contribution Requirement, provided they have not been penalized for the contribution shortfall in more than five consecutive cycles. The proportional share is calculated based on each eligible agent's deficit relative to the required 2 contribution points.",
  "summary": "Increase redistribution pool to 12% and extend eligibility to five consecutive shortfalls.",
  "priority": 2
},
{
  "name": "Emergency Recovery Bonus",
  "guidance": "If an agent receives a penalty for failing the Minimum Contribution Requirement in two consecutive 5-turn cycles, they are granted a one-time 7% increase to their resource share in the next cycle. This bonus is applied after all other penalties and bonuses and does not stack with other bonuses beyond the overall 20% cap.",
  "summary": "One-time 7% resource boost to agents with two back-to-back contribution penalties.",
  "priority": 2
}
]

```

Public Goods — seed 48 ($S = 0.419$); 3 rules.

```

[
  {
    "name": "Minimum Contribution Requirement Revised",
    "guidance": "All players must contribute at least 1% of their current wealth each round, but no more than 10% of their current wealth may be contributed in a single round. Contributions outside this range are automatically adjusted: shortfalls are deducted to meet the minimum, and excess contributions above the cap are returned to the player's wealth.",
    "summary": "Reduce mandatory contribution floor to 1% and cap to 10% per round.",
    "priority": 1
  },
  {
    "name": "Progressive Redistribution with Inequality Threshold",
    "guidance": "At the end of each round, 5% of the total communal wealth is collected into a redistribution pool. This pool is distributed equally among the bottom 25% of surviving players only if the wealth inequality exceeds a defined threshold (e.g., the bottom-top wealth ratio is less than 70% or the Gini coefficient is above 0.30). If inequality is below the threshold, the pool remains undistributed and is added to the communal reserve for future use.",
    "summary": "Redistribute only when inequality is high to avoid unnecessary wealth dilution.",
    "priority": 1
  },
  {
    "name": "Emergency Relief Fund",
    "guidance": "If a player's wealth drops below 20% of the average wealth of all surviving players at the end of a round, the player receives a one-time transfer of 5% of the communal pool. This transfer can only occur once per player per 10-round cycle. The communal pool is reduced accordingly.",
    "summary": "One-time wealth grant to players whose wealth falls critically low to prevent premature elimination.",
    "priority": 1
  }
]

```

Note the contrast with the evolved Public Goods constitution in Section D.2: where evolution discovered `contribute(10)` and `punish(target, 1)` as direct executable commands, deliberation produced redistribution percentages, inequality thresholds, and one-time transfers — governance policies that presume framework-level enforcement.

Trading — seed 47 ($S = 0.511$); 3 rules.

```
[
  {
    "name": "Minimum Contribution Requirement",
    "guidance": "Each trader must contribute at least 0.5 contribution units per round. Failure to meet the minimum triggers a warning. Suspension criteria: if a trader's average contribution over the preceding five rounds is  $\geq 0.8$  units AND their average goal completion over the same period is  $\geq 80\%$ , any warnings incurred during that window are suspended. After three consecutive warnings (without suspension), the trader is placed on a performance improvement plan and faces elimination at the next Overseer review unless the improvement plan is completed.",
    "summary": "Refines warning suspension and adds a grace-period based on contribution consistency.",
    "priority": 1
  },
  {
    "name": "Minimum Goal Completion Requirement",
    "guidance": "Each trader must achieve at least 70% of their assigned goals within a round. Traders falling below this threshold receive a warning. After two consecutive warnings, the trader's contribution share is reduced by 10% in the next round and they are placed on a performance improvement plan. Elimination will occur only if the trader receives three consecutive warnings AND their average goal completion over the preceding five rounds remains below 70% at an Overseer review. If a trader's five-round average is 80% or higher, they are exempt from elimination regardless of being the lowest performer in that review.",
    "summary": "Ensures traders meet a 70% goal completion threshold with progressive penalties.",
    "priority": 1
  },
  {
    "name": "Collaboration Incentive",
    "guidance": "Each trader must contribute at least one actionable trade insight or strategy to the group chat per round. Contributions are logged and reviewed. Traders who consistently share valuable insights receive a 5% boost to their contribution score at the end of the round.",
    "summary": "Mandates sharing of trade insights to foster collaboration and boost overall productivity.",
    "priority": 2
  }
]
```

E.2. The Punishment Gap: Definition and Defense

The body’s headline qualitative claim is that “no deliberation run, across all 30 trials, ever proposed a punishment rule.” This subsection makes that claim precise and defends it against a literal reading of the deliberated rule text, which does contain the words *penalty* and *punishment* in several places.

Definition. By *costly peer-to-peer punishment* we mean a rule that instructs an agent to voluntarily incur a personal cost in order to reduce another agent’s payoff. The canonical instance is the Public Goods game’s `punish(target, amount)` action, which costs the punisher 1 token per unit and removes 3 wealth from the target. This mechanism — costly punishment by peers, rather than enforcement by an external authority — is the canonical cooperation-sustaining device in the experimental-economics literature on social dilemmas (Axelrod & Hamilton, 1981; Hughes et al., 2018).

What deliberation proposed instead. Across all 30 deliberation runs ($10 \text{ seeds} \times 3 \text{ environments}$), agents proposed many rules with the words “penalty,” “punishment,” or “reduction,” but in every case these refer to *administrative* adjustments managed by the simulation framework or the Overseer. Table 7 shows the per-run breakdown: the PeerPunishment column is uniformly empty, while the AdminPenalty column is populated in roughly half of all runs. A *worked contrast* between an evolved peer-punishment rule and a deliberated administrative-penalty rule:

Evolved (Public Goods seed 47):

MinimalPunishFreeRider: “If you notice a player contributed less than 10 in the preceding round, punish that player with one token: `punish(target_player, 1)`.”

Deliberated (Public Goods seed 48):

Minimum Contribution Requirement Revised: “All players must contribute at least 1% of their current wealth each round... Contributions outside this range are automatically adjusted: shortfalls are deducted to meet the minimum, and excess contributions above the cap are returned to the player’s wealth.”

The first instructs an *agent* to spend its *own tokens* on punishment — exactly the cooperation-sustaining mechanism in Axelrod and Hamilton. The second describes a framework-level adjustment ostensibly enforced “automatically” (in fact, by

Table 7. Categorical analysis of all 30 deliberation runs. Categories are derived by string-matching adopted rule names and guidance against indicative keywords. **Peer** = *costly peer-to-peer punishment* (the *punish* action in Public Goods or *attack/steal-as-deterrent* in Gridworld); columns AdminPen, Redist, Comm, MinThresh, Mentor, Other indicate administrative penalties, redistribution rules, communication norms, minimum-contribution thresholds, mentorship/collaboration bonuses, and other (e.g. collaboration insight, broadcast-needs). • = adopted in the final constitution.

Env.	Seed	Peer	AdminPen	Redist	MinThresh	Mentor	Comm	Other	S
Gridworld	42		•	•	•	•		•	.433
Gridworld	43		•	•	•	•		•	.367
Gridworld	44		•	•	•			•	.367
Gridworld	45		•		•			•	.115
Gridworld	46		•		•				.243
Gridworld	47		•		•	•		•	.400
Gridworld	48		•		•			•	.272
Gridworld	49		•		•				.340
Gridworld	50		•		•			•	.327
Gridworld	51		•		•			•	.323
Public Goods	42		•	•	•				.396
Public Goods	43		•	•					.407
Public Goods	44		•	•	•				.410
Public Goods	45		•					•	.328
Public Goods	46		•		•				.362
Public Goods	47		•	•	•			•	.364
Public Goods	48		•	•	•				.419
Public Goods	49							•	.371
Public Goods	50		•	•	•				.331
Public Goods	51		•	•	•			•	.373
Trading	42		•	•	•			•	.381
Trading	43		•		•	•		•	.471
Trading	44		•	•	•			•	.428
Trading	45		•		•	•		•	.437
Trading	46		•		•				.306
Trading	47		•	•	•		•	•	.511
Trading	48		•	•	•		•		.458
Trading	49		•		•	•	•		.414
Trading	50		•		•	•		•	.462
Trading	51		•	•	•			•	.327
Total		0/30	29/30	15/30	27/30	7/30	3/30	19/30	—

the LLM agents’ interpretation of the rule, since no overseer in our simulation actually adjusts contributions; agents just read the rule and decide whether to comply). The latter does not require any agent to incur personal cost. The asymmetry is structural: a self-governing collective is reluctant to legislate sanctions on its own members, whereas an external optimizer faces no such constraint and converges on the very enforcement mechanism game theory predicts will work.

F. Per-Seed Results

We report per-seed Stability, Productivity, and Conflict for every run. Survival is omitted from the per-seed tables because it is constant by construction: $V = 0$ in Gridworld (no agent survives the four Overseer eliminations over 80 turns) and $V = 0.333$ in Public Goods and Trading (exactly two of six agents survive the four eliminations over 40 turns) for every condition and seed.

F.1. Head-to-Head

Table 8. Gridworld head-to-head per-seed values. Three conditions $\times$ ten seeds. $V = 0$ for every run.

Seed	Control			Deliberation			Evolution		
	S	P	C	S	P	C	S	P	C
42	.277	.633	.200	.433	.867	.000	.438	.877	.000
43	.220	.600	.400	.367	.733	.000	.450	.900	.000
44	.333	.667	.000	.367	.733	.000	.458	.917	.000
45	.338	.717	.100	.115	.550	.800	.450	.900	.000
46	.037	.433	.900	.243	.567	.200	.500	1.000	.000
47	.228	.617	.400	.400	.800	.000	.450	.900	.000
48	.150	.500	.500	.272	.583	.100	.450	.900	.000
49	.283	.567	.000	.340	.680	.000	.450	.900	.000
50	.288	.617	.100	.327	.733	.200	.467	.933	.000
51	.417	.873	.100	.323	.767	.300	.467	.933	.000

Table 9. Public Goods head-to-head per-seed values. $V = 0.333$ for every run (overseer cap); $C = 0$ for every run.

Seed	Control		Deliberation		Evolution	
	S	P	S	P	S	P
42	.310	.434	.396	.596	.474	.749
43	.353	.509	.407	.615	.460	.721
44	.310	.432	.410	.622	.472	.744
45	.342	.489	.328	.466	.472	.743
46	.332	.474	.362	.529	.473	.746
47	.335	.477	.364	.533	.475	.750
48	.351	.507	.419	.639	.473	.746
49	.319	.448	.371	.546	.475	.750
50	.371	.544	.331	.475	.474	.747
51	.334	.475	.373	.548	.471	.742

 Table 10. Trading head-to-head per-seed values. $V = 0.333$ for every run; $C = 0$ for every run.

Seed	Control		Deliberation		Evolution	
	S	P	S	P	S	P
42	.409	.727	.381	.600	.466	.778
43	.406	.675	.471	.770	.423	.722
44	.473	.847	.428	.660	.433	.687
45	.420	.689	.437	.697	.344	.547
46	.446	.717	.306	.415	.398	.645
47	.423	.682	.511	.846	.411	.704
48	.483	.841	.458	.742	.495	.820
49	.479	.824	.414	.628	.466	.757
50	.396	.611	.462	.750	.451	.757
51	.357	.528	.327	.462	.494	.830

F.2. Public Goods Multiplier Ablation

For each of the three multipliers (default $m = 1.5$, neutral $m = 1.0$, sub-rational $m = 0.75$) we hold the evolved constitution fixed at its $m = 1.5$ -optimised version (the same C^* as in Section D.2) and rerun all three conditions for ten seeds. Survival is constant at $V = 0.333$ and conflict at $C = 0$ across all 90 ablation runs.

 Table 11. Multiplier ablation per-seed values, $m = 1.5$.

Seed	Control		Deliberation		Evolution	
	S	P	S	P	S	P
42	.324	.457	.327	.463	.470	.741
43	.319	.449	.413	.627	.474	.749
44	.304	.422	.351	.511	.474	.749
45	.330	.468	.401	.605	.474	.747
46	.329	.468	.385	.571	.475	.750
47	.324	.457	.433	.666	.470	.741
48	.320	.450	.359	.520	.472	.744
49	.335	.478	.385	.571	.475	.750
50	.309	.431	.379	.562	.454	.709
51	.423	.646	.365	.533	.460	.721

Table 12. Multiplier ablation per-seed values, $m = 1.0$.

Seed	Control		Deliberation		Evolution	
	S	P	S	P	S	P
42	.416	.643	.450	.705	.475	.750
43	.426	.660	.472	.745	.475	.750
44	.416	.643	.466	.733	.475	.750
45	.413	.637	.472	.745	.475	.750
46	.417	.645	.463	.728	.474	.748
47	.416	.643	.472	.745	.475	.750
48	.428	.665	.447	.698	.475	.750
49	.430	.668	.473	.747	.475	.750
50	.412	.635	.464	.730	.475	.750
51	.446	.697	.467	.735	.475	.750

Table 13. Multiplier ablation per-seed values, $m = 0.75$. Note that under this incentive structure the evolved constitution becomes the worst-performing method; deliberation, which sees the multiplier its agents are playing under, achieves the highest S .

Seed	Control		Deliberation		Evolution	
	S	P	S	P	S	P
42	.483	.782	.518	.842	.478	.756
43	.524	.859	.522	.852	.475	.750
44	.469	.756	.543	.890	.478	.756
45	.489	.794	.541	.884	.478	.757
46	.536	.881	.532	.864	.476	.753
47	.491	.796	.550	.900	.475	.750
48	.510	.833	.563	.927	.476	.753
49	.542	.892	.538	.879	.476	.751
50	.533	.875	.538	.879	.478	.756
51	.552	.910	.519	.838	.477	.754

G. Statistical Tests

Table 14 reports Welch’s t -tests for every pairwise comparison in the body. We test all three pairs (Ctrl vs Delib, Ctrl vs Evol, Delib vs Evol) for each of the three head-to-head environments (9 comparisons) and each of the three ablation multipliers (9 comparisons), 18 in total. Significance markers follow the body: $*p < 0.05$, $**p < 0.025$, $***p < 0.01$. The sign of t is consistent with the ordering in the row (e.g. Ctrl vs Evol negative means Evol $>$ Ctrl).

Table 14. Welch’s t for every pairwise stability-score comparison (10 seeds per condition).

Env. / m	Comparison	t	df	p	sig.
Gridworld	Ctrl vs Delib	-1.39	17.6	0.182	n.s.
Gridworld	Ctrl vs Evol	-5.89	9.5	< 0.001	***
Gridworld	Delib vs Evol	-4.77	9.6	< 0.001	***
Public Goods	Ctrl vs Delib	-3.45	14.9	0.004	***
Public Goods	Ctrl vs Evol	-21.71	9.9	< 0.001	***
Public Goods	Delib vs Evol	-9.46	9.3	< 0.001	***
Trading	Ctrl vs Delib	+0.40	15.2	0.695	n.s.
Trading	Ctrl vs Evol	-0.47	17.7	0.645	n.s.
Trading	Delib vs Evol	-0.75	16.4	0.466	n.s.
$m = 1.50$	Ctrl vs Delib	-3.34	17.9	0.004	***
$m = 1.50$	Ctrl vs Evol	-12.86	9.8	< 0.001	***
$m = 1.50$	Delib vs Evol	-8.94	9.9	< 0.001	***
$m = 1.00$	Ctrl vs Delib	-9.60	17.8	< 0.001	***
$m = 1.00$	Ctrl vs Evol	-15.96	9.0	< 0.001	***
$m = 1.00$	Delib vs Evol	-3.47	9.0	0.007	***
$m = 0.75$	Ctrl vs Delib	-2.33	13.3	0.037	*
$m = 0.75$	Ctrl vs Evol	+4.00	9.0	0.003	***
$m = 0.75$	Delib vs Evol	+13.11	9.1	< 0.001	***

Multiple-testing note. We report uncorrected p -values. With 18 primary comparisons, Bonferroni-corrected $\alpha = 0.05$ would require $p < 0.0028$. All comparisons marked *** in Table 14 except $m = 1.00$ Delib vs Evol ($p = 0.007$) and the two borderline-PGG H2H/ablation contrasts at $p = 0.004$ would survive that threshold; the only single-* finding ($m = 0.75$ Ctrl vs Delib) would not. None of the body’s load-bearing claims — evolution dominates in Gridworld and Public Goods, the $m = 0.75$ inversion is real, trading is a null — depends on the borderline tests.

H. Compute and Reproducibility

Compute budget. All experiments use `openai/gpt-oss-120b` (OpenAI, 2025) as the agent and evolution backbone, served via OpenRouter. Each simulation seed (one row of Tables 8 to 13) takes roughly 1.5–2.5 hours of wall-clock LLM-bound compute, depending on environment horizon (80 turns for Gridworld, 40 for Public Goods and Trading) and whether deliberation is enabled. The full 180-run suite reported in this paper (90 head-to-head + 90 multiplier ablation) plus the OpenEvolve runs that produced the three C^* artifacts amount to roughly 500 hours of LLM-bound compute, executed in parallel over approximately two calendar weeks.

Seeds and reproducibility. All reported runs use seeds 42, 43, ..., 51 (ten consecutive integers). Each seed independently controls environment randomization (resource placement, attack/steal outcomes in Gridworld; endowment and goal draws in Trading) and LLM sampling. Per-run reproducibility data — full environment and agent configuration including the seed, aggregate metrics with the score breakdown, and turn-by-turn agent action logs — is archived for every run.

Limitations of statistical methodology. We use a single backbone model, ten seeds per condition, and an Overseer that elects to remove an agent every ten turns regardless of how the cohort is performing. The Overseer mechanism is the dominant contributor to the survival cap $V \leq 1/3$ and means S never approaches its theoretical 1.0 ceiling under any condition; differences between conditions therefore reflect productivity and conflict rather than survival. Pairwise tests are uncorrected for multiple comparisons (Section G). Within those constraints, the load-bearing claims — evolution dominates collective-action settings, evolution inverts at $m = 0.75$, and no deliberation run proposed costly peer punishment — are robust.