

HEXMACHINA: SELF-EVOLVING MULTI-AGENT SYSTEM FOR CONTINUAL LEARNING OF CATAN

Anonymous authors

Paper under double-blind review

ABSTRACT

We aim to improve on the long-horizon gaps in large language model (LLM) agents by enabling them to sustain coherent strategies in adversarial, stochastic environments. Settlers of Catan provides a challenging benchmark: strategic success depends on balancing short- and long-term goals in the face of dice randomness, trading, expansion, and blocking. This is difficult because prompt-centric LLM agents (e.g., ReAct, Reflexion) must re-interpret large, evolving game states every turn, quickly saturating context windows and failing to maintain consistent strategy across episodes. We propose *HexMachina*, a continual learning multi-agent system that separates environment *discovery* (inducing an adapter layer without documentation) from strategy *improvement* (evolving a compiled player). This architecture preserves executable artifacts, letting the LLM focus on high-level strategy design rather than per-turn decision-making. In controlled Catanatron experiments, HexMachina learns from scratch, evolving players that outperform the strongest human-crafted baseline (AlphaBeta). Our best runs achieve a 54% win rate against AlphaBeta, outperforming prompt-driven LLM agents and shallow no-discovery baselines. Ablations further confirm that greater focus on pure strategy improves performance. Theoretically, this shows that artifact-centric continual learning can transform LLMs from brittle per-turn deciders into stable strategy designers, providing a reusable path toward long-horizon autonomy.

1 INTRODUCTION

Prompt-centric LLM agents and multi-agent systems are powerful, but they struggle on long-horizon tasks: as episodes unfold, prompts saturate with state summaries and ad-hoc "memory," forcing the model to re-interpret the environment at every step (Aghzal et al. (2025); Nayak et al. (2025); Chen et al. (2024)). To move toward autonomous task following and long-horizon competence, an agent should not have to relearn the interface to their environment at each inference step (Bubeck et al. (2023); Park et al. (2023)). This has motivated experimentation with continual learning agent designs that embed feedback loops and let LLMs revise their own prompts and even generate tools/code to improve over time (Zelikman et al. (2022)). In particular, letting an agent gather and preserve artifacts (e.g., reusable functions and typed helpers) offloads heavy context parsing to deterministic code so the model can focus on designing strategy, not re-describing the world.

Despite progress in continual learning, there are few benchmarks that test whether agents can refine a coherent strategy over long horizons. Most existing domains emphasize short tasks or broad skill discovery, offering limited insight into how well an agent can sustain and improve a single competitive policy. Yet this ability is crucial: real-world applications often require agents not just to explore or act locally, but to commit to strategies that hold over many steps in the presence of uncertainty and competition. A benchmark that demands persistent strategy refinement against a strong adversary is therefore essential for evaluating whether lifelong agents truly overcome the long-horizon gap.

Settlers of Catan is an ideal stress test: each turn presents a large, evolving state and action space; success depends on balancing short- and long-term rewards under stochastic resource production, trading, expansion, and adversarial play. Using the open-source Catanatron framework (Collazo (2025)) gives us a controlled interface to observe how a lifelong architecture impacts performance in a domain that reliably exposes limits in long-horizon reasoning.

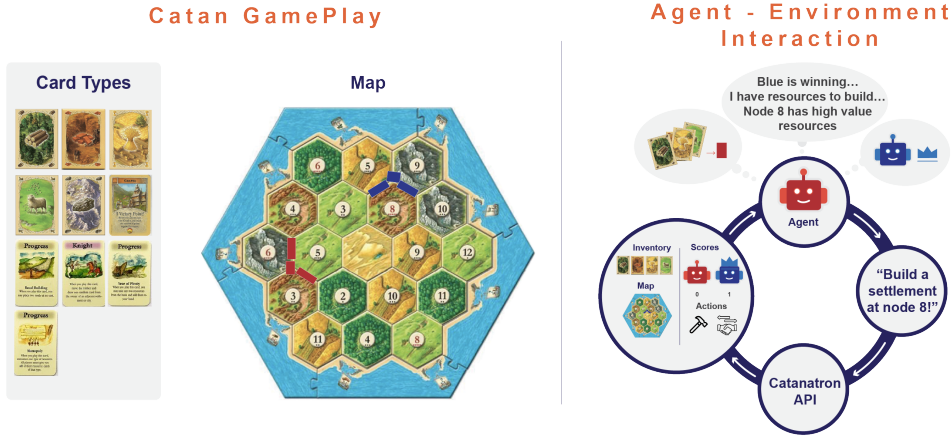


Figure 1: **Overview of Catan gameplay and LLM-agent interaction.** **Left:** *Settlers of Catan* – Players take turns to gather, trade, and spend resources to build on a modular board in a stochastic, partially observable strategy game. The objective is to reach 10 victory points by constructing settlements, roads, and cities Catan Fusion; Catan Collector. **Right:** Our LLM-based framework interacts with the Catanatron API, leveraging game state information and strategic reasoning to decide actions. Through repeated play and self-modification, agents evolve more coherent long-term strategies (Smashicons; murmur (a;b); Hilmy Abiyyu A.; yaicon).

We first demonstrate that traditional per-turn LLM agents (e.g., ReAct/Reflexion-style) perform poorly against a strong human-crafted bot. Asking the model to parse the full game state and independently choose every action while attempting to "hold" a global plan proves unreliable and inconsistent (Table 2). To address this, we separate the act of *thinking* from the act of *playing*, drawing inspiration from the AutoGPT framework (Yang et al. (2023)) to define distinct agent roles: Orchestrator, Analyst, Strategist, Researcher, and Coder. In this configuration, the system hypothesizes a strategy, translates it into a player implementation, reviews the API to ensure correctness, and then evaluates and improves through repeated play. While this Voyager-style (Wang et al. (2023a)) continual learner shows progress, it tends to converge on shallow heuristics that fail to capture the depth of strategic play required in Catan (Appendix A.2). Motivated by this limitation, we introduce a clean separation between the discovery of executable API artifacts and the refinement of strategies built on top of them. With this split, our system, *HexMachina*, evolves players that consistently execute intelligent, long-horizon strategies, outperforming traditional LLM agents, common continual learning architectures, and even the AlphaBeta baseline.

Main Contributions. We highlight the following key contributions from our work:

- **HexMachina: Self-Evolving LLM Agent Framework.** An autonomous system that learns an unknown environment without formal documentation, preserves key code/knowledge as artifacts, and improves its strategy via a closed-loop process that generates and executes code with no human intervention.
- **A strong benchmark setting for continual LLM-agent learning: *Settlers of Catan*.** An environment that both requires long-horizon strategy and distracts naive agents with a large, changing state/action space and delayed rewards.
- **Lifelong agents beat traditional LLM agents on Catan.** HexMachina outperforms prompt-driven baselines and rivals the best human-engineered Catanatron bot (AlphaBeta) by letting the LLM design strategy while compiled code executes it consistently.
- **Empirical importance of separating discovery and improvement.** We show that decoupling environment-artifact discovery from strategy refinement materially improves strategy quality and game performance.

Table 1: **Focus comparison.** ✓=yes, ~=partial, ×=no. Policy evolution (broad: direct or via reward/program/skill search); Artifacts (persisted executable code/skills); Induction (doc-free adapter induction; Voyager ~ with provided control primitives); Adversarial strategy-based (head-to-head vs strong fixed opponent; ✓ only for HexMachina).

System	Environment	Induction	Artifacts	Adversary	Evolution
Voyager	Minecraft	~	✓	×	✓
AlphaEvolve	Code	×	✓	×	✓
Eureka	Isaac Gym	×	✓	×	✓
HexMachina	Catanatron	✓	✓	✓	✓

2 RELATED WORKS

Game-Playing AI and Strategy Games Games have long served as benchmarks for AI research (Gallotta et al. (2024); Costarelli et al. (2024); Nasir et al. (2024)). While significant progress has been made in perfect-information games like Chess and Go (Schultz et al. (2024); Silver et al. (2016)), strategic board games such as *Settlers of Catan*, *Diplomacy* (FAIR) or *Civilization* (Qi et al. (2024)) introduce elements of expanding action spaces, partial observability, and multi-agent interaction, posing unique challenges to an AI system (Szita et al. (2009)). Previous works approached Catan using a specialized neural network architecture to handle its mixed data types, enabling an RL agent to outperform traditional rule-based bots (Gendre & Kaneko (2020)). In contrast, our approach leverages LLMs’ natural language understanding to navigate Catan’s complexities, focusing on autonomous game-play discovery and strategy refinement without relying on extensive training data.

LLM Agents and Long-Horizon Planning LLMs reason well locally but falter at multi-step autonomy: studies report low success on end-to-end plan generation, with models performing better as advisors to external planners Valmeekam et al. (2023). Benchmarks like TravelPlanner confirm poor pass rates even with tools and staged prompting, revealing brittleness under constraint-heavy, multi-objective tasks Xie et al. (2024); Zheng et al. (2025); Nayak et al. (2025); Cui et al. (2025). Prompt-centric agents (ReAct, Reflexion) still act per-turn from ever-growing text context, and multi-agent scaffolds (CAMEL, AutoGen) coordinate via dialogue Yao et al. (2023); Shinn et al. (2023); Wei et al. (2023); Xi et al. (2025); Li et al. (2023); Wu et al. (2023); yet in long-horizon, adversarial domains they repeatedly re-parse large states and lack a persistent executable substrate to enforce strategy across an episode, leaving the planning gap largely intact.

Self-Improvement and Continual Learning Agents Inference-time self-improvement spans verbal reflection (Reflexion), evolutionary prompt search (PromptBreeder, PromptAgent), and code-writing agents that iteratively refine programs (Shinn et al. (2023); Fernando et al. (2023); Wang et al. (2023b)). Surveys systematize these inference-time strategies and the broader landscape of LLM agents (Song et al. (2024); Dong et al. (2024)). Eureka (Ma et al. (2024)) explores program and reward evolution, demonstrating how automated search over reinforcement learning environments can uncover novel control strategies. AlphaEvolve (Novikov et al. (2025)) presents an evolutionary coding agent to tackle open scientific problems and algorithm improvement. Embodied lifelong systems like Voyager show that storing executable skills (a skill library) improves persistence and reuse across episodes, but emphasize breadth (discovering many primitives) rather than depth (refining a single competitive policy).

Building on Voyager, Eureka, and AlphaEvolve, which respectively advance skill discovery, reward/program evolution, and automated code improvement, we shift focus to a different question: can a lifelong LLM system, operating without documentation, induce a compact adapter to an unknown environment and persist executable artifacts in order to evolve a single competitive policy that outperforms traditional LLM agents in adversarial play?

3 BACKGROUND

Settlers of Catan as a Strategic Benchmark *Settlers of Catan* is a 3-4 player board game where players collect and trade resources to build settlements and roads, racing to earn 10 victory points on a modular island map. The game emphasizes resource management, planning, and negotiation, with mechanics like the robber (which blocks resources) adding tactical depth. Catan is known for its balance of luck and skill. **Victory** goes to the first player to reach **10 points**, earned by building and upgrading settlements into cities, buying development cards, and achieving goals like the longest road or largest army. Each settlement is worth 1 point, each city 2, and some development cards grant hidden points or knight bonuses. **Every turn** starts with a dice roll that produces resources for players with adjacent settlements. The active player may then trade and build. If a 7 is rolled, the robber is activated, blocking a tile and stealing a resource. Players must plan expansions, balance upgrades, and trade strategically to manage luck. This need for adaptation and foresight makes Catan a strong benchmark for evaluating strategic reasoning in agents.

The Catanatron Framework We use the open-source Python-based simulator **Catanatron** as our evaluation environment. Designed for automated gameplay of *Settlers of Catan*, Catanatron offers a programmatic interface for integrating custom agents and supports rapid simulation at scale. It faithfully implements the game’s rules and dynamics, capturing key strategic elements such as resource management, trade negotiation via structured proposals, and randomness introduced by dice rolls. Each game consists of players competing to reach ten victory points, with players interacting through well-defined game states that include current resources, board positions, available actions, and observable opponent statuses. Games typically span 40 to 100 turns, allowing for extended observation of agents’ long-term planning capabilities. We benchmark our LLM-driven agents against **AlphaBeta**, the best-performing heuristic agent provided through the API which uses a depth-2 alpha-beta pruning algorithm with heuristic evaluation to select actions.

Alpha-Beta Benchmark Our primary baseline is Catanatron’s AlphaBeta agent: an alpha-beta minimax over stochastic outcomes that computes the expected value of successor states via chance expansion and a fast heuristic value function. Concretely, it uses a depth-2 search (default), a 20 s decision cap, and an optional action-space pruning mode (e.g., robber and maritime-trade pruning heuristics). At leaves, it applies a parameterized value function, and it short-circuits when only one legal action exists. We adopt the author defaults unless otherwise noted, fixing depth = 2 for all reported comparisons. This baseline is both strong and extremely fast, enabling thousands of head-to-head evaluations needed by our continual-learning setup.

4 HEXMACHINA

HexMachina is an autonomous self-evolving multi-agent system that crafts a powerful Catanatron player capable of rivaling the top human-crafted baselines. We utilized Langchain for the model agnostic services, and Langraph for the state machine. Once launched, HexMachina begins by running a *discovery phase*, where it gathers information about the Catanatron API to evolve an **adapters** file. After completion, it enters an *improvement phase* where it begins evolving a **player** file. Each evolution consists of agent collaboration until the Coder writes improvements in the form of testable code. Each phase is limited to 20 evolutions, counted by each time the Coder is called.

4.1 CAPABILITIES

Listed below are the capabilities that enable HexMachina to employ continual learning effectively:

Player Generation HexMachina incrementally codes a complete Catanatron **player** module during the *improvement phase*. The process begins with a minimal template that simply returns the first legal action, then evolves into increasingly sophisticated strategies as feedback accumulates. Importantly, the generated player is not just a script of next actions but an executable policy that can consistently carry out a long-term plan across an entire game. This design shifts the LLM’s role from being a per-turn decider to being a strategy architect, with the Coder agent ensuring that every idea is grounded in syntactically valid and testable code.

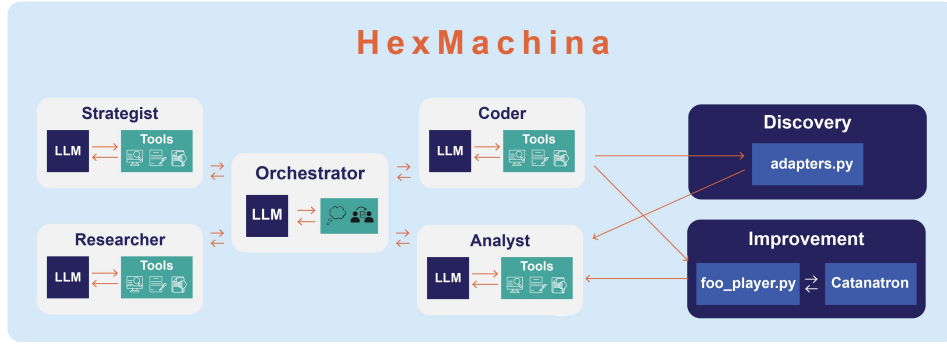


Figure 2: **HexMachina Architecture.** During the discovery phase, the Orchestrator coordinates agents to induce executable functions into the Adapter file, stabilizing access to the environment. In the improvement phase, we found it most effective to rely on a streamlined loop of *Analyst*, *Coder*, and *Orchestrator*, avoiding dilution from additional roles. Here, the Analyst diagnoses performance, the Coder translates revisions into executable code, and the Orchestrator manages iteration. This separation enables the system to refine FooPlayer into a consistent long-horizon strategy.

Experimentation Engine At the core of HexMachina is a deterministic experimentation harness. This engine repeatedly pits evolving players against the strong AlphaBeta baseline under fixed seeds and identical settings, logging outcomes, intermediate states, and decision traces. By holding the environment constant, we can attribute changes in win rate or victory points directly to code evolution, rather than stochastic noise. This repeatable evaluation cycle transforms raw self-play into structured experimental evidence for policy improvement.

Evaluation Evaluation closes the loop: after each batch of games, HexMachina analyzes outcomes to identify which strategic choices were beneficial and which led to failure. This feedback is distilled into concise summaries that the Orchestrator uses to decide whether to preserve, modify, or discard a candidate player. In practice, we found that the most effective evaluation loop did not require every agent’s perspective; instead, a streamlined pipeline of *Analyst to Coder to Orchestrator* yielded clearer strategic signals. Additional recommendations from other roles often diluted the strategy, fragmenting the LLM’s ability to commit to a coherent plan.

Strategy and Discovery HexMachina supports two complementary modes. In the *discovery phase*, it induces executable artifacts such as an `adapters.py` file that stabilizes access to the Catanatron API without any human documentation. This ensures that later improvements build on a reliable, reusable interface. In the *improvement phase*, the system searches over new tactics, revisits prior players, and integrates insights from past runs. Together, these phases ensure that learning is both grounded in the environment and continuously refined across evolutions.

Orchestration The orchestrator serves as the global planner, deciding when to analyze results, request new code, or revisit prior knowledge. Autonomy is enforced by a closed loop: the orchestrator makes high-level decisions based on game outcomes, artifacts, and agent communication, then delegates low-level tasks to the Analyzer and Coder. This separation prevents the system from stalling on details while still maintaining tight control over long-term strategy evolution.

Memory Finally, HexMachina maintains both *game memory* and *semantic memory* across evolutions. Game memory archives past players, their code, and evaluation artifacts, enabling direct comparisons and reuse of successful strategies. Semantic memory allows each agent to persist expertise relevant to its role, e.g., the Coder retaining knowledge of syntax patterns or the Analyst preserving diagnostic heuristics. This dual memory system underpins continual learning: instead of starting from scratch each evolution, the system accumulates strategic and technical knowledge that compounds over time.

4.2 AGENTS

Each agent in HexMachina is a specialist that can call tools (up to 5 per turn) and then yield a single, compact message back to the main loop. Only this final message is persisted to memory, ensuring concise, role-specific contributions. Each agent has access to a *Think Tool* (adapted from Langchain’s deep research agent) to support internal reasoning and explainability. Inputs and outputs are standardized, enabling interchangeable models and providers.

Importantly, not all agents are equally useful in every phase. In the *discovery phase*, the full set of agents contributes to inducing a stable `adapters.py` file from scratch. However, in the *improvement phase*, we found that the most effective configuration is a streamlined loop of **Orchestrator**, **Analyst**, and **Coder**. Additional recommendations from the Strategist and Researcher often diluted coherence, so these roles are used only during discovery or when revisiting artifacts, not for direct strategy refinement.

Orchestrator: Global planner and orchestrator.

Inputs: Orchestrator Messages and summary of evolution

Outputs: Thoughts, system goal, next chosen agent, next agent objective

Tools: None

Coder: Turn strategies into compilable code.

Inputs: Objective from Orchestrator, adapter contents

Outputs: Executable code, summary of changes

Tools: Write/edit file

Analyst: Experimentation evaluator

Inputs: Objective from Orchestrator, summary of evolution, current player and Coder summary of changes, game artifacts, adapter contents

Outputs: Post-game diagnosis, specific analysis, adapter failure

Tools: Read local file

Researcher: Recover API/engine facts and domain tactics. Primarily active during discovery.

Inputs: Objective from Orchestrator, list of files, adapter contents

Outputs: Citations, code pointers, or concise notes with source references

Tools: Read local file, web search

Strategist: Propose concrete, testable plans. Primarily active during discovery.

Inputs: Orchestrator Objective, Evolution Summary, current player, adapter contents

Outputs: Strategy spec and evaluation

Tools: Read local file, view older experiment, web search

5 EXPERIMENT SETUP

We evaluate HexMachina in the open-source *Catanatron* environment under controlled 2-player, 10-point Catan games. Each experiment consists of repeated head-to-head matches against the strongest built-in heuristic bot, *AlphaBeta*. We measure both *win rate* and *final victory points* as indicators of strategic quality. Games are deterministic given a random seed, allowing us to reproduce results and separate genuine improvements from stochastic variance. Data was collected over 60 hours across two machines (MacBook Pro 2019, 16GB; MacBook M1 Max 2021, 32GB).

5.1 BASELINES

Our baselines capture a spectrum of reference points, from trivial random play to a strong, hand-engineered heuristic, allowing us to contextualize HexMachina’s performance against both naive policies and established rule-based expertise.

Random. The simplest control agent chooses uniformly from the legal action space each turn. While strategically meaningless, this baseline sets a lower bound for performance and highlights how much structure even a minimal policy adds.

LLM Player. We also evaluate a Reflexion-style agent (Shinn et al. (2023)) that reformats the game state into text and queries Claude 3.7 once per turn with a high-level goal. This baseline reflects the "prompt-centric" paradigm: the LLM directly drives play without any compiled memory or artifact reuse. Due to inference cost (approx. 70 queries per game), we limited this evaluation to 20 games with a model we had free access too, but it provides a critical comparison to show how quickly context saturation and lack of persistence hinder long-horizon play.

Basic Continual Learner (HexMachina w/o discovery). To isolate the value of separating discovery and improvement, we also test a single-phase continual learning setup equivalent to HexMachina without the discovery phase. Here the system attempts to learn both the environment interface and the strategy simultaneously. This resembles prior lifelong agents such as Voyager and Eureka (Wang et al. (2023a); Ma et al. (2024)), which evolve strategies directly from raw interaction. As shown later in Appendix A.2, these agents often converge on shallow heuristics (e.g., one-ply VP-only evaluators), highlighting the difficulty of strategic refinement without first stabilizing the interface.

AlphaBeta. Finally, we include Catanatron’s AlphaBeta agent, a depth-2 minimax with stochastic expansion and heuristic evaluation. This player is fast, strong, and widely used as a benchmark; in self-play it achieves a 50% win rate by construction. It represents the ceiling for our experiments, providing a human-engineered reference against which HexMachina’s evolved players can be meaningfully compared.

5.2 MODELS

HexMachina is model-agnostic, but in practice we deploy different LLMs for different roles to balance strength and efficiency. We test three orchestrator backends, GPT-5-mini, Claude 3.7, and Mistral-large, to assess robustness across providers. Unless otherwise noted, GPT-5-mini is used for the Coder, which requires reliable code synthesis, while Mistral-large is assigned to support roles (Analyst, Strategist, and Researcher) to reduce cost and latency. This division reflects a general principle of our framework: leverage stronger models where precision is critical (e.g., code generation) and more efficient models where interpretive or diagnostic reasoning suffices.

6 RESULTS AND DISCUSSION

6.1 CONTINUAL LEARNING

We first examine the impact of continual learning through evolution runs of 10 steps, with each step evaluating FooPlayer across 30 games. Figure 3 shows HexMachina steadily improving against AlphaBeta, eventually achieving parity and surpassing baseline players. A central design choice was the separation of *discovery* (API induction and artifact stabilization) from *improvement* (strategy evolution). Our experiments confirm that this separation is critical: systems without discovery struggled to stabilize player code, while those with discovery reliably produced executable players that improved across evolutions.

Interestingly, we found that HexMachina performed better when the Strategist and Researcher agents were *removed*, leaving only the Orchestrator, Analyst, and Coder. While the Strategist was intended to propose concrete plans, results suggest that LLMs often formulate effective strategies in a single shot, and passing these through multiple roles may dilute coherence. Thus, we report results using this streamlined configuration. This insight highlights a broader implication for continual learning: modular multi-agent systems are powerful, but not all roles contribute equally, and reducing mediation can strengthen strategic consistency.

Figure 4 provides a qualitative example of evolution in action. We observe HexMachina iteratively proposing, coding, and refining player strategies while preserving functional artifacts. This illustrates how artifact-centric continual learning transforms an LLM from a per-turn decision maker into a higher-level strategy designer with consistent policy execution.

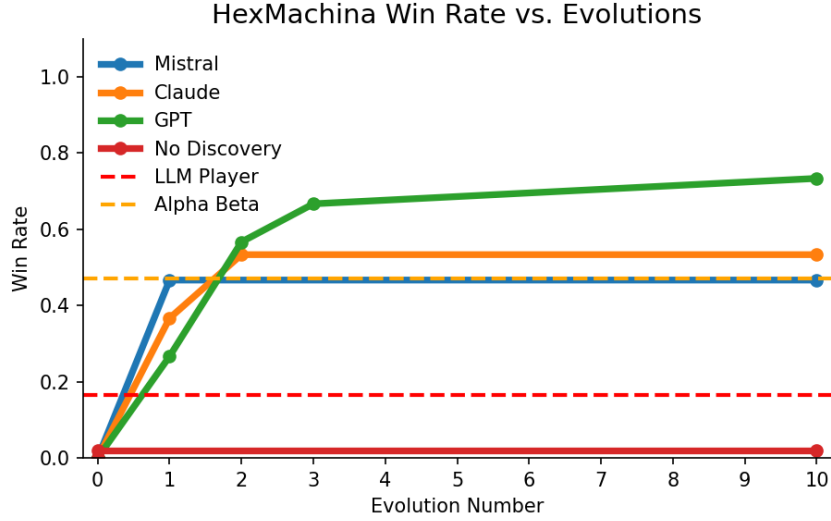


Figure 3: HexMachina Evolving to Outperform Existing Players

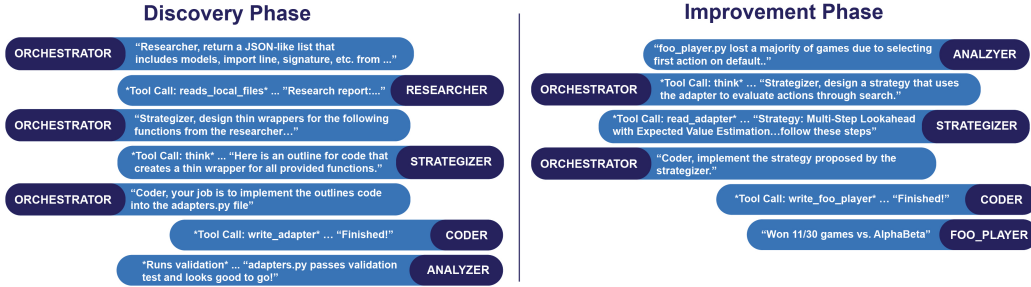


Figure 4: Evolution Messages Example Dialogue

6.2 PLAYER COMPARISON

To stress test the best-evolved players, we ran each configuration 10 times with 100 games per run. Results are summarized in Table 2. HexMachina’s best model (GPT-5-mini) reached a 54.1% win rate and 8.2 ± 0.1 victory points, matching or slightly exceeding AlphaBeta’s 51.0% win rate and 7.8 ± 0.2 points. By contrast, the no-discovery baseline plateaued at much lower win rates, producing players that often failed to generalize beyond static heuristics.

A representative no-discovery agent (Appendix A.2) highlights why this baseline performs poorly. It carries out only a 1-ply lookahead, scoring states almost entirely on current victory points with trivial tie-breakers such as settlements, cities, or roads. With rollouts disabled and no modeling of stochastic production or opponent actions, it assigns identical scores to materially different choices, leading to random tie-breaking, poor settlement placement, and ineffective robber usage. These flaws explain its consistently weak performance in Table 2.

By contrast, the best evolved FooPlayer (Appendix A.1) demonstrates the benefits of HexMachina’s discovery-improvement split. This agent combines phase-aware priorities (early expansion, mid-game balance, late-game upgrades), explicit heuristics for production diversity and robber disruption, and shallow rollouts that anticipate near-term outcomes. These capabilities yield stronger growth, better-timed upgrades, and consistent disruptive pressure on the opponent. The qualitative differences map directly onto the quantitative results in Table 2, underscoring that discovery is critical for stabilizing adapters and enabling the emergence of richer strategies.

Table 2: Win Rate and Victory Points for HexMachina compared to Baselines

Player	Model	Win Rate	Victory Points
HexMachina	GPT	54.1% [51%, 57%]	8.2 $\pm$ 0.1
	Mistral	49.2% [46%, 52%]	7.8 $\pm$ 0.2
	Claude	38.4% [35%, 41%]	7.2 $\pm$ 0.2
LLM Player	Claude	16.4% [3%, 30%]	5.2 $\pm$ 1.2
Alpha-Beta	X	51.0% [48%, 54%]	7.8 $\pm$ 0.2
Random	X	0.2% [0%, 0%]	2.4 $\pm$ 0.0

6.3 ABLATIONS

To isolate the importance of individual design choices, we conducted ablation studies with three independent runs of 10 steps, each tested on 30 games. Results are shown in Table 3. Interestingly, removing the Strategist and Researcher improved performance relative to the full system, reaffirming our earlier finding that direct orchestration leads to clearer strategy translation. Removing the Analyst heavily impacted success as the agent is required to diagnose issues. There would often be situations where the system failed to recognize when functions were being mis-referenced from `adapters.py` without the Analyst bringing it into a failure loop. Overall, these findings back our

Table 3: Multi-Agent Architecture Ablations for HexMachina Policy Evolution

Ablation	Win Rate	Victory Points
All Agents	49.7%	8.0
No Analyst	0.0%	2.1
No Strategist + Researcher	54.1%	8.2

contribution statements: (1) HexMachina evolves executable strategies that rival top human-crafted baselines; (2) artifact preservation and doc-free discovery are essential to this success; (3) LLMs are best deployed at the level of strategy design, not per-turn play; and (4) multi-agent modularity is powerful, but optimal performance may emerge from leaner configurations that avoid unnecessary role handoffs.

7 CONCLUSION

Despite strong results, several limitations hindered performance from improving further. First, we evaluated players solely with win rate and final victory points, coarse metrics that sometimes mask subtler strengths and weaknesses. Second, the LLM occasionally hallucinated code or heuristics, requiring additional filtering. Third, the system was expensive to run due to inference costs, restricting the number of trials. Finally, performance remained closely tied to the quality of the underlying model, with more capable backends producing stronger players. Even with these constraints, HexMachina was able to autonomously induce an API, evolve a robust player, and achieve parity with AlphaBeta, the strongest human-crafted bot.

Looking forward, we see several avenues for advancement. Other researchers could attempt to design a more powerful multi-agent system on this benchmark or build a stronger hand-crafted player for comparison. More broadly, HexMachina should be tested on continual learning benchmarks beyond Catan to validate generality. Finally, the current 20-step evolution limit could be extended with improved memory and player management, enabling longer training horizons and more sophisticated strategies. Together, these extensions would push LLM agents closer to reliable long-horizon autonomy.

8 ETHICAL STATEMENT

Our system executes code in a closed loop with strict safeguards: generated programs run only within a controlled evaluation harness, preventing arbitrary system access. All experiments were logged with fixed random seeds and configuration files, ensuring transparency and reproducibility. While we present HexMachina as an autonomous agent, we avoid anthropomorphizing—our system is a code-evolving tool, not a sentient entity.

9 REPRODUCIBILITY STATEMENT

We release all code, experiment harnesses, and configuration files alongside this submission. To reproduce our results, clone the repository, install dependencies, and follow the step-by-step README instructions. Running experiments requires API keys for the tested LLMs; once provided, the system can be executed exactly as described to replicate all tables and figures.

REFERENCES

- Mohamed Aghzal, Erion Plaku, and Ziyu Yao. Can large language models be good path planners? a benchmark and investigation on spatial-temporal reasoning, 2025. URL <https://arxiv.org/abs/2310.03249>.
- Sébastien Bubeck, Varun Chandrasekaran, Ronen Eldan, Johannes Gehrke, Eric Horvitz, Ece Kamar, Peter Lee, Yin Tat Lee, Yuanzhi Li, Scott Lundberg, Harsha Nori, Hamid Palangi, Marco Tulio Ribeiro, and Yi Zhang. Sparks of artificial general intelligence: Early experiments with gpt-4, 2023. URL <https://arxiv.org/abs/2303.12712>.
- Catan Collector. How to Identify Your Version of Catan. <https://catancollector.com/catan-links/how-to-identify-your-version-of-catan>. Accessed: 2025-05-15.
- Catan Fusion. 3 Royal Weddings in Catan. <http://catanfusion.com/index.php/blog/entry-3-royal-weddings-in-catan>. Accessed: 2025-05-15.
- Yanan Chen, Ali Pesaraghader, Tanmana Sadhu, and Dong Hoon Yi. Can we rely on llm agents to draft long-horizon plans? let’s take travelplanner as an example. *arXiv preprint arXiv:2408.06318*, 2024.
- B. Collazo. Catanatron: Settlers of catan bot simulator and strong ai player. <https://github.com/bcollazo/catanatron>, 2025.
- Anthony Costarelli, Mat Allen, Roman Hauksson, Grace Sodunke, Suhas Hariharan, Carlson Cheng, Wenjie Li, Joshua Clymer, and Arjun Yadav. Gamebench: Evaluating strategic reasoning abilities of llm agents. *arXiv preprint arXiv:2406.06613*, 2024.
- Sijia Cui, Shuai Xu, Aiyao He, Yanna Wang, and Bo Xu. Empowering llms with parameterized skills for adversarial long-horizon planning, 2025. URL <https://arxiv.org/abs/2509.13127>.
- Xiangjue Dong, Maria Teleki, and James Caverlee. A survey on llm inference-time self-improvement. *arXiv preprint arXiv:2412.14352*, 2024.
- Meta Fundamental AI Research Diplomacy Team (FAIR)—Anton Bakhtin, Noam Brown, Emily Dinan, Gabriele Farina, Colin Flaherty, Daniel Fried, Andrew Goff, Jonathan Gray, Hengyuan Hu, et al. Human-level play in the game of diplomacy by combining language models with strategic reasoning. *Science*, 378(6624):1067–1074, 2022.
- Chrisantha Fernando, Dylan Banarse, Henryk Michalewski, Simon Osindero, and Tim Rocktäschel. Promptbreeder: Self-referential self-improvement via prompt evolution. *arXiv preprint arXiv:2309.16797*, 2023.
- Roberto Gallotta, Graham Todd, Marvin Zammit, Sam Earle, Antonios Liapis, Julian Togelius, and Georgios N Yannakakis. Large language models and games: A survey and roadmap. *IEEE Transactions on Games*, 2024.

- Quentin Gendre and Tomoyuki Kaneko. Playing catan with cross-dimensional neural network. In *Neural Information Processing: 27th International Conference, ICONIP 2020, Bangkok, Thailand, November 23–27, 2020, Proceedings, Part II* 27, pp. 580–592. Springer, 2020.
- Hilmy Abiyyu A. Robot icons created by Hilmy Abiyyu A. - Flaticon. <https://www.flaticon.com/free-icons/robot>. Accessed: 2025-05-15.
- Guohao Li, Hasan Abed Al Kader Hammoud, Hani Itani, Dmitrii Khizbullin, and Bernard Ghanem. Camel: Communicative agents for "mind" exploration of large language model society, 2023. URL <https://arxiv.org/abs/2303.17760>.
- Yecheng Jason Ma, William Liang, Guanzhi Wang, De-An Huang, Osbert Bastani, Dinesh Jayaraman, Yuke Zhu, Linxi Fan, and Anima Anandkumar. Eureka: Human-level reward design via coding large language models, 2024. URL <https://arxiv.org/abs/2310.12931>.
- murmur. Conversation icons created by murmur - Flaticon. <https://www.flaticon.com/free-icons/conversation>, a. Accessed: 2025-05-15.
- murmur. Thought bubble icons created by murmur - Flaticon. <https://www.flaticon.com/free-icons/thought-bubble>, b. Accessed: 2025-05-15.
- Muhammad Umair Nasir, Steven James, and Julian Togelius. Gametraversalbenchmark: Evaluating planning abilities of large language models through traversing 2d game maps. *arXiv preprint arXiv:2410.07765*, 2024.
- Siddharth Nayak, Adelmo Morrison Orozco, Marina Ten Have, Vittal Thirumalai, Jackson Zhang, Darren Chen, Aditya Kapoor, Eric Robinson, Karthik Gopalakrishnan, James Harrison, Brian Ichter, Anuj Mahajan, and Hamsa Balakrishnan. Llamar: Long-horizon planning for multi-agent robots in partially observable environments, 2025. URL <https://arxiv.org/abs/2407.10031>.
- Alexander Novikov, Ng     V    , Marvin Eisenberger, Emilien Dupont, Po-Sen Huang, Adam Zolt Wagner, Sergey Shirobokov, Borislav Kozlovskii, Francisco J. R. Ruiz, Abbas Mehrabian, M. Pawan Kumar, Abigail See, Swarat Chaudhuri, George Holland, Alex Davies, Sebastian Nowozin, Pushmeet Kohli, and Matej Balog. Alphaevolve: A coding agent for scientific and algorithmic discovery. Technical report, Google DeepMind, 2025. URL https://colab.research.google.com/github/google-deepmind/alphaevolve_results/blob/master/mathematical_results.ipynb. White paper.
- Joon Sung Park, Joseph C. O’Brien, Carrie J. Cai, Meredith Ringel Morris, Percy Liang, and Michael S. Bernstein. Generative agents: Interactive simulacra of human behavior, 2023. URL <https://arxiv.org/abs/2304.03442>.
- Siyuan Qi, Shuo Chen, Yexin Li, Xiangyu Kong, Junqi Wang, Bangcheng Yang, Pring Wong, Yifan Zhong, Xiaoyuan Zhang, Zhaowei Zhang, et al. Civrealm: A learning and reasoning odyssey in civilization for decision-making agents. *arXiv preprint arXiv:2401.10568*, 2024.
- John Schultz, Jakub Adamek, Matej Jusup, Marc Lanctot, Michael Kaisers, Sarah Perrin, Daniel Hennes, Jeremy Shar, Cannada Lewis, Anian Ruoss, et al. Mastering board games by external and internal planning with language models. *arXiv preprint arXiv:2412.12119*, 2024.
- Noah Shinn, Federico Cassano, Edward Berman, Ashwin Gopinath, Karthik Narasimhan, and Shunyu Yao. Reflexion: Language agents with verbal reinforcement learning, 2023. URL <https://arxiv.org/abs/2303.11366>.
- David Silver, Aja Huang, Chris J. Maddison, Arthur Guez, Laurent Sifre, George van den Driessche, Julian Schrittwieser, Ioannis Antonoglou, Vedavyas Panneshelvam, Marc Lanctot, Sander Dieleman, Dominik Grewe, John Nham, Nal Kalchbrenner, Ilya Sutskever, Timothy P. Lillicrap, Madeleine Leach, Koray Kavukcuoglu, Thore Graepel, and Demis Hassabis. Mastering the game of go with deep neural networks and tree search. *Nature*, 529(7587):484–489, 2016. doi: 10.1038/nature16961.
- Smashicons. Trade icons created by Smashicons - Flaticon. <https://www.flaticon.com/free-icons/trade>. Accessed: 2025-05-15.

- Yuda Song, Hanlin Zhang, Carson Eisenach, Sham Kakade, Dean Foster, and Udaya Ghai. Mind the gap: Examining the self-improvement capabilities of large language models. *arXiv preprint arXiv:2412.02674*, 2024.
- István Szita, Guillaume Chaslot, and Pieter Spronck. Monte-carlo tree search in settlers of catan. In *Advances in computer games*, pp. 21–32. Springer, 2009.
- Karthik Valmeekam, Sarath Sreedharan, Matthew Marquez, Alberto Olmo, and Subbarao Kambhampati. On the planning abilities of large language models (a critical investigation with a proposed benchmark), 2023. URL <https://arxiv.org/abs/2302.06706>.
- Guanzhi Wang, Yuqi Xie, Yunfan Jiang, Ajay Mandlekar, Chaowei Xiao, Yuke Zhu, Linxi Fan, and Anima Anandkumar. Voyager: An open-ended embodied agent with large language models, 2023a. URL <https://arxiv.org/abs/2305.16291>.
- Xinyuan Wang, Chenxi Li, Zhen Wang, Fan Bai, Haotian Luo, Jiayou Zhang, Nebojsa Jojic, Eric P Xing, and Zhiting Hu. Promptagent: Strategic planning with language models enables expert-level prompt optimization. *arXiv preprint arXiv:2310.16427*, 2023b.
- Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Brian Ichter, Fei Xia, Ed Chi, Quoc Le, and Denny Zhou. Chain-of-thought prompting elicits reasoning in large language models, 2023. URL <https://arxiv.org/abs/2201.11903>.
- Qingyun Wu, Gagan Bansal, Jieyu Zhang, Yiran Wu, Beibin Li, Erkang Zhu, Li Jiang, Xiaoyun Zhang, Shaokun Zhang, Jiale Liu, Ahmed Hassan Awadallah, Ryen W White, Doug Burger, and Chi Wang. Autogen: Enabling next-gen llm applications via multi-agent conversation, 2023. URL <https://arxiv.org/abs/2308.08155>.
- Zhiheng Xi, Wenxiang Chen, Xin Guo, Wei He, Yiwen Ding, Boyang Hong, Ming Zhang, Junzhe Wang, Senjie Jin, Enyu Zhou, et al. The rise and potential of large language model based agents: A survey. *Science China Information Sciences*, 68(2):121101, 2025.
- Jian Xie, Kai Zhang, Jiangjie Chen, Tinghui Zhu, Renze Lou, Yuandong Tian, Yanghua Xiao, and Yu Su. Travelplanner: A benchmark for real-world planning with language agents, 2024. URL <https://arxiv.org/abs/2402.01622>.
- yaicon. Build icons created by yaicon - Flaticon. <https://www.flaticon.com/free-icons/build>. Accessed: 2025-05-15.
- Hui Yang, Sifu Yue, and Yunzhong He. Auto-gpt for online decision making: Benchmarks and additional opinions, 2023. URL <https://arxiv.org/abs/2306.02224>.
- Shunyu Yao, Jeffrey Zhao, Dian Yu, Nan Du, Izhak Shafran, Karthik Narasimhan, and Yuan Cao. React: Synergizing reasoning and acting in language models. In *International Conference on Learning Representations (ICLR)*, 2023.
- Eric Zelikman, Yuhuai Wu, Jesse Mu, and Noah D. Goodman. Star: Bootstrapping reasoning with reasoning, 2022. URL <https://arxiv.org/abs/2203.14465>.
- Junhao Zheng, Chengming Shi, Xidi Cai, Qiuke Li, Duzhen Zhang, Chenxing Li, Dong Yu, and Qianli Ma. Lifelong learning of large language model based agents: A roadmap, 2025. URL <https://arxiv.org/abs/2501.07278>.

A APPENDIX

A.1 HEXMACHINA’S BEST STRATEGY

```
import random
from typing import Iterable, List, Optional, Any, Tuple

# MUST use the adapters surface to interact with the game environment
from .adapters import (
```

```

648     Game,
649     Player,
650     Color,
651     copy_game,
652     execute_deterministic,
653     execute_spectrum,
654     expand_spectrum,
655     list_prunned_actions,
656     prune_robber_actions,
657     base_fn,
658     value_production,
659     get_value_fn,
660 )
661
662 class FooPlayer(Player):
663     """A Foo player with game-phase aware decisioning, improved sampling,
664     short rollouts, and richer heuristics.
665
666     This implementation is defensive: it uses only the adapters surface
667     and contains many fallbacks when attributes or adapter helpers are
668     missing.
669
670     Key features:
671     - Game-phase detection (early/mid/late) to bias settlement/road vs
672       city/dev-card
673     - Settlement & road potential heuristics to encourage early expansion
674     - Robber/knight evaluation to value disruption and steals
675     - Must-include guarantees for critical action types (settlement/road/
676       robber/dev)
677     - Rollout policy biased by phase and includes a light opponent-
678       response
679
680     NOTE: Many game model attribute names vary across environments. This
681     code attempts multiple common attribute names and falls back to string-
682     based heuristics when necessary. If the next run raises AttributeError for
683     an adapters function or a specific attribute, provide the traceback so
684     it can be patched to the concrete environment.
685     """
686
687     # Tunable constants (exposed to edit for experimentation)
688     MAX_SIMULATIONS = 24
689     PREFILTER_TOP_K = 8
690     ROLLOUT_DEPTH = 2
691     SIMULATION_BUDGET = 60
692     DEBUG = False
693
694     # Phase thresholds (used by get_game_phase)
695     EARLY_TURN_THRESHOLD = 20
696     MID_TURN_THRESHOLD = 45
697
698     # Phase multipliers matrix (explicit)
699     MULTS = {
700         "EARLY": {"settlement": 2.0, "road": 1.8, "city": 0.8, "dev": 1.2},
701         "MID": {"settlement": 1.0, "road": 1.0, "city": 1.25, "dev": 1.0},
702         "LATE": {"settlement": 0.8, "road": 0.9, "city": 1.5, "dev": 1.0},
703     }

```

```

702
703 # Must-include action tokens (robust, lowercase matching)
704 MUST_INCLUDE_TOKENS = {
705     "build_city",
706     "build_settlement",
707     "build_sett",
708     "build_road",
709     "buy_dev",
710     "buy_dev_card",
711     "buycard",
712     "play_knight",
713     "knight",
714     "move_robber",
715     "move_robber_action",
716     "robber",
717     "trade",
718     "offer_trade",
719 }
720
721 # Robber scoring base (increased)
722 ROBBER_BASE_SCORE = 80.0
723 ROBBER_BASE_SCORE_HIGH = 80.0
724
725 # Settlement target in early game
726 TARGET_SETTLEMENTS_EARLY = 3
727
728 # Epsilon-greedy randomness to avoid predictability
729 EPSILON_GREEDY = 0.04
730
731 # Rollout bonuses for the very first rollout step
732 ROLLOUT_SETTLEMENT_BONUS = 1.7
733 ROLLOUT_ROAD_BONUS = 1.4
734
735 # Tie tolerance
736 TOLERANCE = 1e-6
737
738 # Development card deck & EV constants
739 DEV_DECK = {"knight": 14, "vp": 5, "road_building": 2, "
740             year_of_plenty": 2, "monopoly":
741             2}
742
743 DEV_TOTAL = sum(DEV_DECK.values())
744 EV_KNIGHT = 0.15
745 EV_VP = 1.0
746 EV_ROAD_BUILDING = 0.25
747 EV_YOP = 0.2
748 EV_MONOPOLY = 0.3
749 DEV_EV_SCALE = 60.0
750 DEV_EV_THRESHOLD = 0.25
751
752 # Knight bonuses
753 KNIGHT_LARGEST_ARMY_BONUS = 50.0
754 KNIGHT_BASE = 25.0
755 KNIGHT_MIN_SCORE = 35.0
756
757 # City/road/robber tuning (from latest analyzer guidance)
758 CITY_URGENCY_BONUS = 85.0
759 CITY_AFFORD_STRICT_ORE = 3
760 CITY_AFFORD_STRICT_WHEAT = 2
761 CITY_AFFORD_SOON_ORE = 2
762 CITY_AFFORD_SOON_WHEAT = 1
763 ROLLOUT_CITY_BONUS = 1.8
764 ROAD_SCORE_BOOST = 9.0
765 PROD_LOSS_IMPORTANCE = 70.0
766 HIGH_VALUE_RESOURCE_SET = {"ore", "wheat", "metal", "grain"}
767 CITY_TIE_EPS = 0.02

```

```

756
757 # Forcing behavior flags and diagnostic counters
758 PREFILTER_FORCE_CITY_IF = True
759 CITY_FORCE_AFFORD_STRICT = True
760 DEBUG_COUNTS = False
761
762 def __init__(self, name: Optional[str] = None):
763     super().__init__(Color.BLUE, name)
764     # Try to cache a base value function from adapters
765     try:
766         self._value_fn = base_fn()
767         self.debug_print("FooPlayer: Using adapters.base_fn() for
768                             evaluation")
769     except Exception as e:
770         self._value_fn = None
771         self.debug_print("FooPlayer: adapters.base_fn() not available
772                             , will use heuristic.
773                             Error:", e)
774
775 # Diagnostic counters (quiet unless DEBUG)
776 self._diag_forced_settlement = 0
777 self._diag_forced_road = 0
778 self._diag_city_urgency_count = 0
779 self._diag_settle_urgency_count = 0
780
781 # New counters for tuning
782 self.COUNTER_FORCED_CITY = 0
783 self.COUNTER_DEV_BUY_FORCED = 0
784 self.COUNTER_BUY_DEV_ACTUALLY = 0
785 self.COUNTER_BUILD_CITY_ACTUALLY = 0
786 self.COUNTER_ROBBER_ACTUALLY = 0
787
788 # ----- Debug helper -----
789 def debug_print(self, *args: Any) -> None:
790     if self.DEBUG:
791         print(*args)
792
793 # ----- Utility helpers -----
794 def _get_player_color(self) -> Color:
795     """Return this player's color. Try common attribute names."""
796     if hasattr(self, "color"):
797         return getattr(self, "color")
798     if hasattr(self, "_color"):
799         return getattr(self, "_color")
800     return Color.BLUE
801
802 def _safe_action_name(self, action: Any) -> str:
803     """Produce a lowercase string name for the action for robust
804         matching."""
805     try:
806         at = getattr(action, "action_type", None)
807         if at is None:
808             at = getattr(action, "type", None)
809         if at is not None:
810             try:
811                 return str(at.name).lower()
812             except Exception:
813                 return str(at).lower()
814     except Exception:
815         pass
816     try:
817         # Some Action objects have a .name or .action_name
818         name = getattr(action, "name", None) or getattr(action, "
819                                                         action_name", None)
820         if name is not None:

```

```

810         return str(name).lower()
811     except Exception:
812         pass
813     try:
814         return str(action).lower()
815     except Exception:
816         return ""
817
818 # ----- Phase detection -----
819 def get_game_phase(self, game: Game, color: Optional[Color] = None) -
820     > str:
821     """Return 'EARLY', 'MID', or 'LATE' based on turn counters or VP
822     thresholds.
823
824     Order of checks:
825     1) turn/tick counters if available (preferred)
826     2) max VP among players
827     3) fallback to conservative MID
828     """
829     try:
830         state = getattr(game, "state", game)
831         turn_count = (
832             getattr(state, "turn", None)
833             or getattr(state, "tick", None)
834             or getattr(state, "turn_count", None)
835             or getattr(state, "tick_count", None)
836         )
837         if isinstance(turn_count, (int, float)):
838             tc = int(turn_count)
839             if tc < self.EARLY_TURN_THRESHOLD:
840                 return "EARLY"
841             if tc < self.MID_TURN_THRESHOLD:
842                 return "MID"
843             return "LATE"
844     except Exception:
845         pass
846
847 # Fallback: use maximum VP among players
848 try:
849     state = getattr(game, "state", game)
850     players = getattr(state, "players", None) or getattr(game, "
851     players", None) or []
852     max_vp = 0
853     if isinstance(players, dict):
854         for p in players.values():
855             vp = getattr(p, "victory_points", None) or getattr(p,
856             "vp", None) or
857             0
858             try:
859                 vp = int(vp)
860             except Exception:
861                 vp = 0
862             max_vp = max(max_vp, vp)
863     else:
864         for p in players:
865             vp = getattr(p, "victory_points", None) or getattr(p,
866             "vp", None) or
867             0
868             try:
869                 vp = int(vp)
870             except Exception:
871                 vp = 0
872             max_vp = max(max_vp, vp)
873     if max_vp < 4:
874         return "EARLY"

```

```

864         if max_vp < 8:
865             return "MID"
866         return "LATE"
867     except Exception:
868         # Conservative fallback to MID
869         return "MID"
870
871 # ----- Heuristic / evaluation (phase-aware) -----
872 def _heuristic_value(self, game: Game, color: Color) -> float:
873     """Phase-aware heuristic including production potential and city-
874         upgrade progress.
875
876     Many attribute names are attempted to be robust across different
877         game models.
878
879     # Die probabilities for numbers 2..12 ignoring 7
880     die_prob = {2: 1 / 36, 3: 2 / 36, 4: 3 / 36, 5: 4 / 36, 6: 5 / 36
881                 , 8: 5 / 36, 9: 4 / 36, 10: 3 / 36, 11: 2 / 36, 12: 1 /
882                 36}
883
884     # Player lookup
885     player_state = None
886     try:
887         state = getattr(game, "state", game)
888         players = getattr(state, "players", None) or getattr(game, "
889             players", None)
890         if isinstance(players, dict):
891             player_state = players.get(color) or players.get(str(
892                 color))
893         elif isinstance(players, (list, tuple)):
894             for p in players:
895                 if getattr(p, "color", None) == color or getattr(p, "
896                     color", None) ==
897                     str(color):
898                     player_state = p
899                     break
900     except Exception:
901         player_state = None
902
903     def _safe_get(obj, *names, default=0):
904         if obj is None:
905             return default
906         for name in names:
907             try:
908                 val = getattr(obj, name)
909                 if val is not None:
910                     return val
911             except Exception:
912                 try:
913                     val = obj[name]
914                     if val is not None:
915                         return val
916                 except Exception:
917                     continue
918         return default
919
920     vp = _safe_get(player_state, "victory_points", "vp", default=0)
921     settlements = _safe_get(player_state, "settlements", "
922         settle_count", "
923         settle_locations", default=0
924     )
925     if isinstance(settlements, (list, tuple)):
926         settlements = len(settlements)

```

```

918 cities = _safe_get(player_state, "cities", "city_count", "
919                                     city_locations", default=0)
920 if isinstance(cities, (list, tuple)):
921     cities = len(cities)
922 roads = _safe_get(player_state, "roads", "road_count", default=0)
923 if isinstance(roads, (list, tuple)):
924     roads = len(roads)
925 dev_vp = _safe_get(player_state, "dev_vp", "dev_victory_points",
926                                     default=0)
927
928 # Resources summary
929 resources_obj = _safe_get(player_state, "resources", default=0)
930 resources_total = 0
931 resource_diversity = 0
932 try:
933     if isinstance(resources_obj, dict):
934         resources_total = sum(resources_obj.values())
935         resource_diversity = sum(1 for v in resources_obj.values()
936                                 if v > 0)
937     elif isinstance(resources_obj, (list, tuple)):
938         resources_total = sum(resources_obj)
939         resource_diversity = sum(1 for v in resources_obj if v >
940                                 0)
941     else:
942         resources_total = int(resources_obj)
943         resource_diversity = 1 if resources_total > 0 else 0
944 except Exception:
945     resources_total = 0
946     resource_diversity = 0
947
948 # Production potential estimation
949 prod_value = 0.0
950 try:
951     board = getattr(state, "board", None) or getattr(game, "board",
952                                                         None)
953     hexes = getattr(board, "hexes", None) or getattr(board, "
954                                                         tiles", None) or []
955     settlements_list = _safe_get(player_state, "settlements", "
956                                     settle_locations",
957                                     default=[])
958     if isinstance(settlements_list, (list, tuple)):
959         for s in settlements_list:
960             try:
961                 for h in hexes:
962                     neighbors = getattr(h, "vertices",
963                                         None) or getattr(h, "adjacent_vertices",
964                                                         None) or []
965                     if s in neighbors:
966                         num = getattr(h, "roll",
967                                       None) or getattr(h, "number",
968                                                         None) or getattr(h, "value",
969                                                         None)
970                         try:
971                             num = int(num)
972                         except Exception:
973                             num = None
974                         if num in die_prob:
975                             prod_value += die_prob[num] * 1.0
976             except Exception:
977                 continue
978     cities_list = _safe_get(player_state, "cities", "
979                                     city_locations", default
980                                     =[])
981     if isinstance(cities_list, (list, tuple)):
982         for c in cities_list:

```

```

972         try:
973             for h in hexes:
974                 neighbors = getattr(h, "vertices",
975                                     None) or
976                                     getattr(h, "adjacent_vertices", None) or []
977                 if c in neighbors:
978                     num = getattr(h, "roll",
979                                   None) or getattr(h, "number",
980                                                       None) or getattr(h, "value",
981                                                                    None)
982                     try:
983                         num = int(num)
984                     except Exception:
985                         num = None
986                     if num in die_prob:
987                         prod_value += die_prob[num] * 2.0
988             except Exception:
989                 continue
990 except Exception:
991     prod_value = 0.0
992
993 # City upgrade progress heuristic
994 city_resource_val = 0.0
995 try:
996     if isinstance(resources_obj, dict):
997         wheat = resources_obj.get("wheat", 0) + resources_obj.get(
998             "grain", 0)
999         ore = resources_obj.get("ore", 0) + resources_obj.get("
1000             metal", 0)
1001         city_resource_val = min(wheat, ore)
1002 except Exception:
1003     city_resource_val = 0.0
1004
1005 # Phase multipliers
1006 phase = self.get_game_phase(game, color)
1007 mults = self.MULTS.get(phase, self.MULTS["MID"])
1008 settlement_mul = mults["settlement"]
1009 road_mul = mults["road"]
1010 city_mul = mults["city"]
1011 dev_mul = mults["dev"]
1012
1013 # Adjust production weight by phase
1014 prod_weight = 80.0 if phase == "EARLY" else 45.0 if phase == "MID
1015                  " else 30.0
1016
1017 # Compose weighted sum (city reward scaled by city_mul)
1018 score = (
1019     float(vp) * 100.0
1020     + float(settlements) * 25.0 * settlement_mul
1021     + float(cities) * 60.0 * city_mul
1022     + float(roads) * 6.0 * road_mul
1023     + float(dev_vp) * 50.0
1024     + float(resources_total) * 1.0
1025     + float(resource_diversity) * 3.0
1026     + float(city_resource_val) * 5.0
1027     + float(prod_value) * prod_weight
1028 )
1029
1030 return float(score)
1031
1032 def _evaluate_game_state(self, game: Game, color: Color) -> float:
1033     """Evaluate a single game state for the given player color.
1034
1035     Prefer adapters.base_fn() if available (cached in self._value_fn)
1036         . If available, combine

```

```

1026         it with the heuristic for stability. We keep phase multipliers
1027             inside the heuristic so
1028         they influence the final blended value.
1029         """
1030         heuristic = self._heuristic_value(game, color)
1031         if self._value_fn is not None:
1032             try:
1033                 vf_val = float(self._value_fn(game, color))
1034                 return 0.85 * vf_val + 0.15 * heuristic
1035             except Exception as e:
1036                 self.debug_print("FooPlayer: value_fn failed during
1037                                     evaluate_game_state,
1038                                     falling back to
1039                                     heuristic. Error:",
1040                                     e)
1041         return float(heuristic)
1042
1043     # ----- Cheap scoring & potentials -----
1044     def _get_player_state(self, game: Game, color: Color) -> Any:
1045         """Return the player_state object from the game state (best-
1046             effort)."""
1047         try:
1048             state = getattr(game, "state", game)
1049             players = getattr(state, "players", None) or getattr(game, "
1050                                     players", None)
1051             if isinstance(players, dict):
1052                 return players.get(color) or players.get(str(color))
1053             elif isinstance(players, (list, tuple)):
1054                 for p in players:
1055                     if getattr(p, "color", None) == color or getattr(p, "
1056                                     color", None) ==
1057                                     str(color):
1058                         return p
1059         except Exception:
1060             return None
1061         return None
1062
1063     def settlement_potential(self, action: Any, game: Game, color: Color)
1064         -> float:
1065         """Estimate benefit of a settlement action: new resource types
1066             and production.
1067
1068         Best-effort: try to parse adjacent hexes from action or fallback
1069             to string heuristics.
1070         """
1071         bonus = 0.0
1072         try:
1073             name = self._safe_action_name(action)
1074             # Quick check: if action indicates a settlement, give base
1075             if any(tok in name for tok in ("build_settlement", "
1076                                     build_sett", "settle")):
1077                 bonus += 5.0
1078
1079             # Try to parse a vertex index from the action string
1080             digits = [int(tok) for tok in name.split() if tok.isdigit()]
1081             vertex = digits[0] if digits else None
1082
1083             state = getattr(game, "state", game)
1084             board = getattr(state, "board", None) or getattr(game, "board
1085                                     ", None)
1086             hexes = getattr(board, "hexes", None) or getattr(board, "
1087                                     tiles", None) or []
1088
1089             # Player's current resource types
1090             player_state = self._get_player_state(game, color)

```

```

1080     player_types = set()
1081     try:
1082         settlements_list = getattr(player_state, "settlements",
1083                                     None) or getattr(
1084                                         player_state, "
1085                                         settle_locations",
1086                                         None) or []
1087         if isinstance(settlements_list, (list, tuple)):
1088             for s in settlements_list:
1089                 for h in hexes:
1090                     neighbors = getattr(h, "vertices",
1091                                         None) or getattr(h, "adjacent_vertices",
1092                                                         None) or []
1093                     if s in neighbors:
1094                         rtype = getattr(h, "resource",
1095                                         None) or getattr(h, "type",
1096                                                         None)
1097                         if rtype is not None:
1098                             player_types.add(str(rtype).lower())
1099     except Exception:
1100         player_types = set()
1101
1102     # Adjacent resources for proposed vertex
1103     adj_resources = set()
1104     prod_sum = 0.0
1105     die_prob = {2: 1 / 36, 3: 2 / 36, 4: 3 / 36, 5: 4 / 36, 6: 5
1106                / 36, 8: 5 / 36, 9: 4 /
1107                36, 10: 3 / 36, 11: 2 /
1108                36, 12: 1 / 36}
1109
1110     if vertex is not None:
1111         for h in hexes:
1112             try:
1113                 neighbors = getattr(h, "vertices",
1114                                     None) or getattr(h, "adjacent_vertices",
1115                                                         None) or []
1116                 if vertex in neighbors:
1117                     rtype = getattr(h, "resource",
1118                                     None) or getattr(h, "type",
1119                                                         None)
1120                     if rtype is not None:
1121                         adj_resources.add(str(rtype).lower())
1122                     num = getattr(h, "roll",
1123                                   None) or getattr(h, "number",
1124                                                       None) or getattr(h, "value",
1125                                                           None)
1126                     try:
1127                         num = int(num)
1128                     except Exception:
1129                         num = None
1130                     if num in die_prob:
1131                         prod_sum += die_prob[num]
1132             except Exception:
1133                 continue
1134
1135     # New types
1136     new_types = adj_resources - player_types
1137     bonus += float(len(new_types)) * 12.0
1138     bonus += float(prod_sum) * 8.0
1139 except Exception:
1140     pass
1141 return float(bonus)
1142
1143 def road_connection_potential(self, action: Any, game: Game, color:
1144                               Color) -> float:
1145     """Estimate if a road action helps expansion. Best-effort using
1146         indices."""

```

```

1134     bonus = 0.0
1135     try:
1136         name = self._safe_action_name(action)
1137         # try to extract numbers from action name
1138         digits = [int(tok) for tok in name.split() if tok.isdigit()]
1139         # player's settlement/city vertices
1140         player_state = self._get_player_state(game, color)
1141         player_nodes = set()
1142         try:
1143             settles = getattr(player_state, "settlements", None) or
1144                         getattr(player_state, "settle_locations",
1145                                 None) or []
1146             cities = getattr(player_state, "cities", None) or getattr(
1147                 player_state, "city_locations",
1148                 None) or []
1149             if isinstance(settles, (list, tuple)):
1150                 player_nodes.update(settles)
1151             if isinstance(cities, (list, tuple)):
1152                 player_nodes.update(cities)
1153         except Exception:
1154             player_nodes = set()
1155
1156         if digits:
1157             # if any digit matches a player node, give higher bonus
1158             if any(d in player_nodes for d in digits):
1159                 bonus += 6.0
1160             else:
1161                 bonus += 3.0
1162         else:
1163             # fallback string heuristics
1164             if "build_road" in name or ("road" in name and "build" in
1165                                         name):
1166                 bonus += 2.0
1167         except Exception:
1168             pass
1169         return float(bonus)
1170
1171     def evaluate_buy_dev_card(self, action: Any, game: Game, color: Color
1172                             ) -> bool:
1173         """Decide whether buying a dev card is currently a good idea (
1174             best-effort)."""
1175         try:
1176             player_state = self._get_player_state(game, color)
1177             resources = getattr(player_state, "resources", None)
1178             if isinstance(resources, dict):
1179                 ore = resources.get("ore", 0) + resources.get("metal", 0)
1180                 wheat = resources.get("wheat", 0) + resources.get("grain",
1181                                                                     0)
1182                 others = sum(v for k, v in resources.items() if k not in
1183                             ("ore", "metal", "wheat", "grain"))
1184                 # if have ore+wheat+another, prefer dev card; or if no
1185                 # settlement/road/city affordable
1186                 if ore >= 1 and wheat >= 1 and others >= 1:
1187                     return True
1188                 # fallback: if early game and we have some resources but
1189                 # no settlement potential, allow dev buy
1190                 phase = self.get_game_phase(game, color)
1191                 if phase == "EARLY" and (ore + wheat + others) >= 3:
1192                     return True

```

```

1188         except Exception:
1189             pass
1190         return False
1191
1192     def dev_card_ev_estimate(self, game: Game, color: Color) -> float:
1193         """Estimate expected VP-equivalent value of buying a development
1194             card.
1195
1196         Uses static DEV_DECK and EV_* constants and scales by opponent
1197             pressure and army gaps.
1198         Returns a small VP-equivalent number (e.g., ~0.3-0.6 when
1199             favorable).
1200         """
1201         try:
1202             base_ev = 0.0
1203             # composition-based base EV
1204             base_ev += (self.DEV_DECK.get("knight", 0) / self.DEV_TOTAL)
1205                 * self.EV_KNIGHT
1206             base_ev += (self.DEV_DECK.get("vp", 0) / self.DEV_TOTAL) *
1207                 self.EV_VP
1208             base_ev += (self.DEV_DECK.get("road_building", 0) / self.
1209                 DEV_TOTAL) * self.
1210                 EV_ROAD_BUILDING
1211             base_ev += (self.DEV_DECK.get("year_of_plenty", 0) / self.
1212                 DEV_TOTAL) * self.EV_YOP
1213             base_ev += (self.DEV_DECK.get("monopoly", 0) / self.DEV_TOTAL
1214                 ) * self.EV_MONOPOLY
1215
1216             # Scale factors: opponents production pressure and army
1217                 proximity
1218             # Compute opponents' max production (best-effort)
1219             state = getattr(game, "state", game)
1220             board = getattr(state, "board", None) or getattr(game, "board",
1221                 None)
1222             hexes = getattr(board, "hexes", None) or getattr(board, "
1223                 tiles", None) or []
1224
1225             opponents = []
1226             players = getattr(state, "players", None) or getattr(game, "
1227                 players", None) or []
1228             my_color = color
1229             if isinstance(players, dict):
1230                 for k, p in players.items():
1231                     if k == my_color or getattr(p, "color", None) ==
1232                         my_color:
1233                         continue
1234                     opponents.append(p)
1235             else:
1236                 for p in players:
1237                     if getattr(p, "color", None) == my_color:
1238                         continue
1239                     opponents.append(p)
1240
1241             # compute simple production score for each opponent
1242             die_prob = {2: 1 / 36, 3: 2 / 36, 4: 3 / 36, 5: 4 / 36, 6: 5
1243                 / 36, 8: 5 / 36, 9: 4 /
1244                 36, 10: 3 / 36, 11: 2 /
1245                 36, 12: 1 / 36}
1246
1247             max_opp_prod = 0.0
1248             for opp in opponents:
1249                 prod = 0.0
1250                 opp_settles = getattr(opp, "settlements", None) or
1251                     getattr(opp, "
1252                     settle_locations",
1253                     None) or []

```

```

1242         opp_cities = getattr(opp, "cities", None) or getattr(opp,
1243                             "city_locations",
1244                             None) or []
1245     try:
1246         for s in opp_settles:
1247             for h in hexes:
1248                 neighbors = getattr(h, "vertices",
1249                                     None) or getattr(h, "adjacent_vertices",
1250                                     None) or []
1251                 if s in neighbors:
1252                     num = getattr(h, "roll",
1253                                   None) or getattr(h, "number",
1254                                   None) or getattr(h, "value",
1255                                   None)
1256                     try:
1257                         num = int(num)
1258                     except Exception:
1259                         num = None
1260                     if num in die_prob:
1261                         prod += die_prob[num]
1262     for c in opp_cities:
1263         for h in hexes:
1264             neighbors = getattr(h, "vertices",
1265                                 None) or getattr(h, "adjacent_vertices",
1266                                 None) or []
1267             if c in neighbors:
1268                 num = getattr(h, "roll",
1269                               None) or getattr(h, "number",
1270                               None) or getattr(h, "value",
1271                               None)
1272                 try:
1273                     num = int(num)
1274                 except Exception:
1275                     num = None
1276                 if num in die_prob:
1277                     prod += 2.0 * die_prob[num]
1278     except Exception:
1279         pass
1280     max_opp_prod = max(max_opp_prod, prod)
1281
1282     # army gap factor
1283     my_state = self._get_player_state(game, color)
1284     my_army = getattr(my_state, "army", None) or getattr(my_state,
1285                                                         "army_size", None) or
1286             getattr(my_state, "knights_played", None)
1287             or 0
1288     try:
1289         my_army = int(my_army)
1290     except Exception:
1291         my_army = 0
1292     max_other_army = 0
1293     try:
1294         if isinstance(players, dict):
1295             for k, p in players.items():
1296                 if k == my_color or getattr(p, "color", None) ==
1297                                         my_color:
1298                     continue
1299                 oa = getattr(p, "army", None) or
1300                     getattr(p, "army_size", None) or
1301                     getattr(p, "knights_played",
1302                             None) or 0
1303                 try:
1304                     oa = int(oa)
1305                 except Exception:

```

```

1296         oa = 0
1297         max_other_army = max(max_other_army, oa)
1298     else:
1299         for p in players:
1300             if getattr(p, "color", None) == my_color:
1301                 continue
1302             oa = getattr(p, "army", None) or
1303                 getattr(p, "army_size", None) or
1304                 getattr(p, "knights_played", None) or 0
1305             try:
1306                 oa = int(oa)
1307             except Exception:
1308                 oa = 0
1309             max_other_army = max(max_other_army, oa)
1310     except Exception:
1311         max_other_army = 0
1312
1313     army_gap = max(0, max_other_army - my_army)
1314
1315     # scale base_ev conservatively
1316     scale = 1.0
1317     if max_opp_prod > 0.25: # opponent has strong production
1318         scale += 0.25
1319     if army_gap >= 1:
1320         scale += 0.15 * army_gap
1321
1322     final_ev = base_ev * scale
1323     return float(final_ev)
1324 except Exception:
1325     # fallback conservative
1326     return 0.25
1327
1328 def build_urgency(self, game: Game, color: Color) -> Tuple[float,
1329                                                             float, float]:
1330     """Return (city_bonus, settlement_bonus, road_bonus) depending on
1331         resources and phase."""
1332
1333     city_bonus = 0.0
1334     settlement_bonus = 0.0
1335     road_bonus = 0.0
1336     try:
1337         player_state = self._get_player_state(game, color)
1338         resources = getattr(player_state, "resources", None) or {}
1339         if not isinstance(resources, dict):
1340             # try to coerce
1341             try:
1342                 total = sum(resources)
1343                 resources = {"res": total}
1344             except Exception:
1345                 resources = {}
1346
1347     # simple can_afford_city_soon heuristic
1348     ore = resources.get("ore", 0) + resources.get("metal", 0)
1349     wheat = resources.get("wheat", 0) + resources.get("grain", 0)
1350     settlements_list = getattr(player_state, "settlements", None)
1351                       or getattr(player_state,
1352                                "settle_locations",
1353                                None) or []
1354     settlements_owned = len(settlements_list) if isinstance(
1355         settlements_list, (list,
1356                           tuple)) else 0
1357
1358     phase = self.get_game_phase(game, color)
1359     # If mid/late and can afford city soon, large city urgency
1360     if phase in ("MID", "LATE") and ore >= 2 and wheat >= 1:
1361         city_bonus += 40.0

```

```

1350         self._diag_city_urgency_count += 1
1351         # If early and lacking settlements target, encourage
1352             settlements strongly
1353         if phase == "EARLY" and settlements_owned < self.
1354             TARGET_SETTLEMENTS_EARLY
1355             :
1356             settlement_bonus += 35.0
1357             self._diag_settle_urgency_count += 1
1358             # Road potential: give moderate constant bonus
1359             road_bonus += 10.0
1359         except Exception:
1360             pass
1361         return city_bonus, settlement_bonus, road_bonus
1362
1363     def cheap_pre_score(self, action: Any, game: Game, color: Color) ->
1364         float:
1365         """Cheap, fast scoring used to prioritize actions for simulation
1366             (phase-aware)."""
1367
1368         s = 0.0
1369         name = self._safe_action_name(action)
1370
1371         phase = self.get_game_phase(game, color)
1372         mults = self.MULTS.get(phase, self.MULTS["MID"])
1373         settlement_mul = mults["settlement"]
1374         road_mul = mults["road"]
1375         city_mul = mults["city"]
1376         dev_mul = mults["dev"]
1377
1378         # urgency bonuses
1379         city_urgency, sett_urgency, road_urgency = self.build_urgency(
1380             game, color)
1381
1382         # Reward direct VP gains but adjust city bias early
1383         if any(tok in name for tok in ("build_city",)):
1384             base_city = max(50.0, 100.0 * city_mul - 15.0)
1385             # penalize city if early and still below settlement target
1386             try:
1387                 player_state = self._get_player_state(game, color)
1388                 settles = getattr(player_state, "settlements", None) or
1389                     getattr(player_state, "settle_locations",
1390                         None) or []
1391                 curr_settlements = len(settles) if isinstance(settles, (
1392                     list, tuple)) else 0
1393                 if phase == "EARLY" and curr_settlements <
1394                     self.TARGET_SETTLEMENTS_EARLY:
1395                     base_city *= 0.6
1396             except Exception:
1397                 pass
1398             s += base_city + city_urgency
1399
1400         if any(tok in name for tok in ("build_settlement", "build_sett"))
1401             :
1402             s += 90.0 * settlement_mul
1403             # add settlement potential (resource diversity / production)
1404             s += self.settlement_potential(action, game, color) * (1.0 if
1405                 phase != "EARLY" else
1406                 settlement_mul)
1407
1408             s += sett_urgency
1409
1410         if "buy_dev" in name or "buycard" in name or "buy_dev_card" in
1411             name:
1412             # compute EV estimate
1413             dev_ev = self.dev_card_ev_estimate(game, color)
1414             s += dev_ev * self.DEV_EV_SCALE

```

```

1404         # slightly reduced base bias to favor cities when urgent
1405         if self.evaluate_buy_dev_card(action, game, color):
1406             s += 8.0 * dev_mul
1407         try:
1408             if dev_ev >= self.DEV_EV_THRESHOLD:
1409                 s += 2.0
1410         except Exception:
1411             pass
1412
1413         if "build_road" in name or ("road" in name and "build" in name):
1414             s += 20.0 * road_mul
1415             s += self.road_connection_potential(action, game, color) * (1
1416                                                         .0 if phase != "EARLY"
1417                                                         else road_mul)
1418             s += road_urgency
1419
1420         if "knight" in name or "play_knight" in name:
1421             # raise baseline and include army/steal bonuses
1422             s += 70.0
1423             s += self.evaluate_play_knight(action, game, color)
1424
1425         if "robber" in name or "move_robber" in name:
1426             s += 50.0
1427             s += self.evaluate_robber_action(action, game, color)
1428
1429         if "trade" in name or "offer_trade" in name:
1430             s += 10.0
1431
1432         # Encourage hitting settlement target early
1433         try:
1434             player_state = self._get_player_state(game, color)
1435             curr_settlements = 0
1436             settles = getattr(player_state, "settlements", None) or
1437                         getattr(player_state, "
1438                             settle_locations", None)
1439                         or []
1440
1441             if isinstance(settles, (list, tuple)):
1442                 curr_settlements = len(settles)
1443             if phase == "EARLY" and curr_settlements < self.
1444                                     TARGET_SETTLEMENTS_EARLY
1445                                     and any(tok in name for
1446                                         tok in ("
1447                                             build_settlement", "
1448                                             build_sett"))):
1449                 s += 30.0
1450         except Exception:
1451             pass
1452
1453         # small settlement/road potentials for other actions
1454         if not any(tok in name for tok in ("build_settlement", "
1455                                             build_sett")):
1456             s += self.settlement_potential(action, game, color) * 0.1
1457         if not any(tok in name for tok in ("build_road",)):
1458             s += self.road_connection_potential(action, game, color) * 0.
1459                                     1
1460
1461         # Minor random tie-break
1462         s += random.random() * 1e-3
1463         return s
1464
1465         # ----- Prefilter actions (phase-aware guarantees)
1466         -----
1467         def prefilter_actions(self, actions: List[Any], game: Game, color:
1468                                 Color) -> List[Any]:

```

```

1458         """Return a bounded list of candidate actions to evaluate
1459             thoroughly.
1460
1461         Guarantees inclusion of must-include tokens and early-game
1462             settlement/road actions.
1463         """
1464         if not actions:
1465             return []
1466
1467         all_actions = list(actions)
1468         phase = self.get_game_phase(game, color)
1469
1470         musts = []
1471         others = []
1472         found_settlement = None
1473         found_road = None
1474         for a in all_actions:
1475             name = self._safe_action_name(a)
1476             if any(tok in name for tok in self.MUST_INCLUDE_TOKENS):
1477                 if a not in musts:
1478                     musts.append(a)
1479             else:
1480                 others.append(a)
1481                 if found_settlement is None and any(tok in name for tok in ("
1482                     build_settlement", "
1483                     build_sett", "settle")):
1484                     found_settlement = a
1485                 if found_road is None and any(tok in name for tok in ("
1486                     build_road", "road")):
1487                     found_road = a
1488
1489         # Phase-based forced includes: ensure at least one settlement and
1490             one road action if present
1491             in EARLY
1492         if phase == "EARLY":
1493             if found_settlement is not None and found_settlement not in
1494                 musts:
1495                 musts.append(found_settlement)
1496                 self._diag_forced_settlement += 1
1497             if found_road is not None and found_road not in musts:
1498                 musts.append(found_road)
1499                 self._diag_forced_road += 1
1500
1501         # Include recommended dev-card buys if conservative and EV
1502             threshold met
1503         for a in all_actions:
1504             name = self._safe_action_name(a)
1505             if any(tok in name for tok in ("buy_dev", "buycard", "
1506                 buy_dev_card")):
1507                 try:
1508                     if self.evaluate_buy_dev_card(a, game, color):
1509                         dev_ev = self.dev_card_ev_estimate(game, color)
1510                         if dev_ev >= self.DEV_EV_THRESHOLD and a not in
1511                             musts:
1512                             # include only if dev EV merits it
1513                             musts.append(a)
1514                 except Exception:
1515                     pass
1516
1517         # Ensure robber/knight actions are present
1518         for a in all_actions:
1519             name = self._safe_action_name(a)
1520             if any(tok in name for tok in ("robber", "move_robber", "
1521                 knight", "play_knight")):
1522                 :

```

```

1512         if a not in musts:
1513             musts.append(a)
1514
1515     # Score and pick top-K from others
1516     scored = [(self.cheap_pre_score(a, game, color), a) for a in
1517               others]
1518     scored.sort(key=lambda x: x[0], reverse=True)
1519     top_k = [a for (_, a) in scored[: self.PREFILTER_TOP_K]]
1520
1521     # Combine unique musts + top_k preserving order
1522     candidates = []
1523     for a in musts + top_k:
1524         if a not in candidates:
1525             candidates.append(a)
1526
1527     # Fill up with random remaining samples until MAX_SIMULATIONS
1528     remaining = [a for a in all_actions if a not in candidates]
1529     random.shuffle(remaining)
1530     while len(candidates) < min(len(all_actions), self.
1531                               MAX_SIMULATIONS) and
1532           remaining:
1533         candidates.append(remaining.pop())
1534
1535     if not candidates and all_actions:
1536         candidates = random.sample(all_actions, min(len(all_actions),
1537                                                     self.MAX_SIMULATIONS))
1538
1539     self.debug_print(f"FooPlayer: Prefilter selected {len(candidates)}
1540                      candidates (musts={len(
1541                      musts)}, phase={phase})")
1542
1543     if self.DEBUG and phase == "EARLY":
1544         self.debug_print(f"    Forced includes: settlement={'yes' if
1545                      found_settlement else '
1546                      no'}, road={'yes' if
1547                      found_road else 'no'}")
1548
1549     return candidates
1550
1551 # ----- Playable actions extraction -----
1552 def get_playable_actions_from_game(self, game: Game) -> List[Any]:
1553     """Try adapters.list_prunned_actions first, then common game
1554     attributes."""
1555
1556     try:
1557         acts = list_prunned_actions(game)
1558         if acts:
1559             return acts
1560     except Exception as e:
1561         self.debug_print("FooPlayer: list_prunned_actions unavailable
1562                          or failed. Error:", e)
1563
1564     try:
1565         if hasattr(game, "get_playable_actions"):
1566             return list(game.get_playable_actions())
1567     except Exception:
1568         pass
1569
1570     try:
1571         if hasattr(game, "playable_actions"):
1572             return list(getattr(game, "playable_actions"))
1573     except Exception:
1574         pass
1575
1576     try:
1577         state = getattr(game, "state", None)
1578         if state is not None and hasattr(state, "playable_actions"):
1579             return list(getattr(state, "playable_actions"))
1580     except Exception:
1581         pass

```

```

1566
1567         return []
1568
1569     # ----- Robber / Knight evaluation -----
1570     def evaluate_robber_action(self, action: Any, game: Game, color:
1571                               Color) -> float:
1572         """Estimate the value of moving the robber (best-effort).
1573
1574         If the action does not specify a target hex, evaluate all hexes
1575         and prefer the
1576         one that maximizes opponent production loss.
1577         """
1578         score = 0.0
1579         try:
1580             # Base preference to include robber moves (use HIGH base for
1581             # aggressive play)
1582             score += self.ROBBER_BASE_SCORE_HIGH
1583             name = self._safe_action_name(action)
1584             # Try to parse a target hex id
1585             digits = [int(tok) for tok in name.split() if tok.isdigit()]
1586             target = digits[0] if digits else None
1587
1588             # Die probabilities
1589             die_prob = {2: 1 / 36, 3: 2 / 36, 4: 3 / 36, 5: 4 / 36, 6: 5
1590                        / 36, 8: 5 / 36, 9: 4 /
1591                        36, 10: 3 / 36, 11: 2 /
1592                        36, 12: 1 / 36}
1593
1594             state = getattr(game, "state", game)
1595             board = getattr(state, "board", None) or getattr(game, "board",
1596                                                             None)
1597             hexes = getattr(board, "hexes", None) or getattr(board, "
1598                                                             tiles", None) or []
1599
1600             # Map hex identifier to object (best-effort: use index or id)
1601             hex_map = {}
1602             for idx, h in enumerate(hexes):
1603                 try:
1604                     hid = getattr(h, "id", None) or getattr(h, "index",
1605                                                             None) or idx
1606                 except Exception:
1607                     hid = idx
1608                 try:
1609                     key = int(hid) if isinstance(hid, int) or (isinstance
1610                                                                (hid, str) and
1611                                                                hid.isdigit())
1612                                                                else idx
1613                 except Exception:
1614                     key = idx
1615                 hex_map[key] = h
1616
1617             # Determine best target if none specified
1618             targets_to_consider = [target] if target in hex_map else list
1619                                     (hex_map.keys())
1620
1621             # Compute production loss on opponents per candidate target
1622             opponents = []
1623             players = getattr(state, "players", None) or getattr(game, "
1624                                                             players", None) or []
1625             my_color = color
1626             if isinstance(players, dict):
1627                 for k, p in players.items():
1628                     if k == my_color or getattr(p, "color", None) ==
1629                                             my_color:
1630                         continue

```

```

1620         opponents.append(p)
1621     else:
1622         for p in players:
1623             if getattr(p, "color", None) == my_color:
1624                 continue
1625             opponents.append(p)
1626
1627     best_loss = 0.0
1628     best_steal = 0.0
1629     best_hex = None
1630     resource_value = {"ore": 3.0, "metal": 3.0, "wheat": 3.0, "
1631                      grain": 3.0, "brick": 2.
1632                      0, "lumber": 2.0, "wood"
1633                      : 2.0, "sheep": 2.0}
1634
1635     for t in targets_to_consider:
1636         try:
1637             if t not in hex_map:
1638                 continue
1639             h = hex_map[t]
1640             num = getattr(h, "roll", None) or getattr(h, "number"
1641              , None) or
1642              getattr(h, "
1643              value", None)
1644
1645             try:
1646                 num = int(num)
1647             except Exception:
1648                 num = None
1649             prob = die_prob.get(num, 0)
1650             total_prod_loss = 0.0
1651             steal_expected = 0.0
1652             for opp in opponents:
1653
1654                 opp_settles = getattr(opp,
1655                  "settlements", None) or getattr(opp,
1656                  "settle_locations", None) or []
1657                 opp_cities = getattr(opp, "cities",
1658                  None) or getattr(opp, "city_locations",
1659                  None) or []
1660                 mult = 0.0
1661                 try:
1662                     for s in opp_settles:
1663                         neighbors = getattr(h,
1664                          "vertices", None) or getattr(h,
1665                          "adjacent_vertices", None) or []
1666                         if s in neighbors:
1667                             mult += 1.0
1668                     for c in opp_cities:
1669                         neighbors = getattr(h,
1670                          "vertices", None) or getattr(h,
1671                          "adjacent_vertices", None) or []
1672                         if c in neighbors:
1673                             mult += 2.0
1674                 except Exception:
1675                     continue
1676             total_prod_loss += prob * mult
1677             # Estimate steal expected
1678             try:
1679                 opp_resources = getattr(opp,
1680                  "resources", None) or {}
1681                 if isinstance(opp_resources, dict)
1682                 and opp_resources:
1683                     total_res =
1684                     sum(opp_resources.values())
1685                     if total_res > 0:

```

```

1674         avg_val =
1675             sum(resource_value.get(r,
1676                 1.5) * (opp_resources.get(r,
1677                     0) / total_res) for r in
1678                 opp_resources)
1679         steal_expected += avg_val *
1680             0.5
1681     except Exception:
1682         pass
1683     # choose best
1684     if total_prod_loss > best_loss or (abs(
1685         total_prod_loss
1686         - best_loss) <
1687         1e-9 and
1688         steal_expected >
1689         best_steal):
1690         best_loss = total_prod_loss
1691         best_steal = steal_expected
1692         best_hex = t
1693     except Exception:
1694         continue
1695
1696     # Aggressive scaling per latest tuning
1697     score += best_loss * self.PROD_LOSS_IMPORTANCE
1698     score += best_steal * 30.0
1699     # Extra bonus if multiple opponent cities affected
1700     try:
1701         if best_hex in hex_map:
1702             h = hex_map[best_hex]
1703             city_count = 0
1704             for opp in opponents:
1705                 for c in getattr(opp, "cities", []) or
1706                     getattr(opp, "city_locations", []) or []:
1707                     neighbors = getattr(h, "vertices",
1708                         None) or getattr(h,
1709                         "adjacent_vertices", None) or []
1710                     if c in neighbors:
1711                         city_count += 1
1712             if city_count > 0:
1713                 score += 20.0 * city_count
1714     except Exception:
1715         pass
1716
1717     # If steal estimated is very significant, add
1718     decisive bonus
1719     if best_steal > 2.0:
1720         score += 30.0
1721
1722     # Debug
1723     if self.DEBUG and best_hex is not None:
1724         self.debug_print(f"FooPlayer: evaluate_robber_action
1725                             best_hex={best_hex}
1726                             prod_loss={best_loss
1727                                 :.3f} steal_ev={
1728                                     best_steal:.2f}")
1729
1730     except Exception:
1731         pass
1732     return float(score)
1733
1734 def evaluate_play_knight(self, action: Any, game: Game, color: Color)
1735     -> float:
1736     """Estimate the value of playing a knight (best-effort)."""
1737     score = float(self.KNIGHT_BASE)
1738     try:

```

```

1728         name = self._safe_action_name(action)
1729         if "steal" in name or "rob" in name:
1730             score += 10.0
1731
1732         # army progress
1733         player_state = self._get_player_state(game, color)
1734         army = getattr(player_state, "army", None) or getattr(
1735             player_state, "army_size", None) or getattr(
1736             player_state, "
1737             knights_played", None)
1738             or 0
1739
1740         try:
1741             army = int(army)
1742         except Exception:
1743             army = 0
1744
1745         # detect largest army threshold
1746         largest_threshold = 3
1747         try:
1748             state = getattr(game, "state", game)
1749             players = getattr(state, "players", None) or getattr(game,
1750                 "players", None)
1751                 or []
1752
1753             max_other = 0
1754             if isinstance(players, dict):
1755                 for k, p in players.items():
1756                     if getattr(p, "color", None) == color or k ==
1757                         color:
1758                         continue
1759                     other_army = getattr(p, "army", None) or getattr(
1760                         p, "
1761                         army_size",
1762                         None) or
1763                         getattr(p, "
1764                         knights_played",
1765                         None) or
1766                         0
1767
1768                     try:
1769                         other_army = int(other_army)
1770                     except Exception:
1771                         other_army = 0
1772                     max_other = max(max_other, other_army)
1773
1774             else:
1775                 for p in players:
1776                     if getattr(p, "color", None) == color:
1777                         continue
1778                     other_army = getattr(p, "army", None) or getattr(
1779                         p, "
1780                         army_size",
1781                         None) or
1782                         getattr(p, "
1783                         knights_played",
1784                         None) or
1785                         0
1786
1787                     try:
1788                         other_army = int(other_army)
1789                     except Exception:
1790                         other_army = 0
1791                     max_other = max(max_other, other_army)
1792             largest_threshold = max(3, max_other + 1)
1793         except Exception:
1794             largest_threshold = 3
1795
1796         if army + 1 >= largest_threshold:

```

```

1782         score += self.KNIGHT_LARGEST_ARMY_BONUS
1783     else:
1784         score += 20.0
1785
1786     # Debug
1787     if self.DEBUG:
1788         self.debug_print(f"FooPlayer: evaluate_play_knight army={
1789                                     army} target={
1790                                     largest_threshold}
1791                                     score={score}")
1792
1793     except Exception:
1794         pass
1795     return float(score)
1796
1797 # ----- Helper: determine active player color
1798 -----
1799 def _get_active_player_color(self, game: Game) -> Optional[Color]:
1800     """Best-effort to detect which Color is to move in the given game
1801     state."""
1802
1803     try:
1804         state = getattr(game, "state", game)
1805         cp = getattr(state, "current_player", None) or getattr(state,
1806                                     "active_player", None)
1807         or getattr(state, "
1808                                     turn_color", None)
1809
1810         if cp is None:
1811             cp = getattr(game, "current_player", None)
1812             # cp might be index, player object, or Color
1813             if isinstance(cp, Color):
1814                 return cp
1815             if isinstance(cp, int):
1816                 players = getattr(state, "players", None) or getattr(game,
1817                                     "players", None)
1818                 or []
1819
1820                 try:
1821                     if isinstance(players, (list, tuple)) and 0 <= cp <
1822                                     len(players):
1823                         return getattr(players[cp], "color", None)
1824                     except Exception:
1825                         pass
1826             # If cp is a player object
1827             if hasattr(cp, "color"):
1828                 return getattr(cp, "color")
1829
1830             # Fallback: pick first player in players whose color != our
1831             color
1832             players = getattr(state, "players", None) or getattr(game, "
1833                                     players", None) or []
1834             my_color = self._get_player_color()
1835             if isinstance(players, dict):
1836                 for k, p in players.items():
1837                     try:
1838                         c = getattr(p, "color", None) or k
1839                         if c != my_color:
1840                             return c
1841                     except Exception:
1842                         continue
1843             else:
1844                 for p in players:
1845                     try:
1846                         c = getattr(p, "color", None)
1847                         if c != my_color:
1848                             return c
1849                     except Exception:
1850                         continue

```

```

1836         except Exception:
1837             pass
1838         return None
1839
1840     # ----- Rollout logic with opponent-response -----
1841     def rollout_value(self, game: Game, color: Color, depth: int, initial
1842                       : bool = True) -> float:
1843         """Short greedy rollout with phase bias and light opponent-
1844           response.
1845
1846           initial: True for the first step of rollout so we can bias toward
1847           expansion early.
1848         """
1849         try:
1850             if depth <= 0:
1851                 return self._evaluate_game_state(game, color)
1852
1853             actions = self.get_playable_actions_from_game(game)
1854             if not actions:
1855                 return self._evaluate_game_state(game, color)
1856
1857             phase = self.get_game_phase(game, color)
1858
1859             def score_for_rollout(a, g, c, is_initial):
1860                 base = self.cheap_pre_score(a, g, c)
1861                 if is_initial and phase == "EARLY":
1862                     name = self._safe_action_name(a)
1863                     if any(tok in name for tok in ("build_settlement", "
1864                                                     build_sett", "
1865                                                     settle")):
1866                         base *= self.ROLLOUT_SETTLEMENT_BONUS
1867                     if any(tok in name for tok in ("build_road", "road")):
1868                         base *= self.ROLLOUT_ROAD_BONUS
1869                 return base
1870
1871             sorted_actions = sorted(actions, key=lambda a:
1872                                   score_for_rollout(a,
1873                                                       game, color, initial),
1874                                   reverse=True)
1875
1876             # Try top actions to simulate
1877             for a in sorted_actions[:6]:
1878                 branches = []
1879                 try:
1880                     branches = execute_deterministic(game, a)
1881                 except Exception:
1882                     try:
1883                         branches = execute_spectrum(game, a)
1884                     except Exception:
1885                         branches = []
1886                 if not branches:
1887                     continue
1888                 # pick the most probable branch
1889                 next_game = max(branches, key=lambda bp: float(bp[1]))[0]
1890
1891                 # Light opponent-response: if opponent to move next,
1892                 # simulate their greedy action once
1893                 opp_color = self._get_active_player_color(next_game)
1894                 my_color = color
1895                 if opp_color is not None and opp_color != my_color and
1896                   depth >= 2:
1897                     try:

```

```

1890         opp_actions = self.get_playable_actions_from_game
1891                             (next_game)
1892         if opp_actions:
1893             # filter out robber/knight for
1894             opponent response unless all are
1895             robber/knight
1896             non_disrupt = [oa for oa in
1897             opp_actions if not any(tok in
1898             self._safe_action_name(oa) for tok
1899             in ("knight", "robber",
1900             "move_robber"))]
1901             candidate_ops = non_disrupt if
1902             non_disrupt else opp_actions
1903             # pick opponent best action by
1904             cheap_pre_score from their
1905             perspective
1906             best_opp = max(candidate_ops,
1907             key=lambda oa:
1908             self.cheap_pre_score(oa, next_game,
1909             opp_color))
1910             # simulate opponent action
1911             deterministically if possible
1912             opp_branches = []
1913             try:
1914                 opp_branches =
1915                 execute_deterministic(next_game,
1916                 best_opp)
1917             except Exception:
1918                 try:
1919                     opp_branches =
1920                     execute_spectrum(next_game,
1921                     best_opp)
1922                 except Exception:
1923                     opp_branches = []
1924             if opp_branches:
1925                 next_game = max(opp_branches,
1926                 key=lambda bp: float(bp[1]))[0]
1927         except Exception:
1928             pass
1929
1930         return self.rollout_value(next_game, color,
1931         depth - 1, initial=False)
1932
1933     # fallback: try any action that simulates
1934     for a in sorted_actions[:10]:
1935         branches = []
1936         try:
1937             branches = execute_deterministic(game, a)
1938         except Exception:
1939             try:
1940                 branches = execute_spectrum(game, a)
1941             except Exception:
1942                 branches = []
1943         if branches:
1944             next_game = max(branches, key=lambda bp:
1945             float(bp[1]))[0]
1946             return self.rollout_value(next_game, color,
1947             depth - 1, initial=False)
1948
1949     return self._evaluate_game_state(game, color)
1950 except Exception as e:
1951     self.debug_print("FooPlayer: rollout_value exception, falling
1952                       back to
1953                       evaluate_game_state.
1954                       Error:", e)

```

```

1944         return self._evaluate_game_state(game, color)
1945
1946     # ----- Evaluate action expectation (enhanced)
1947     -----
1948     def _evaluate_action_expectation(self, game: Game, action: Any,
1949                                     per_action_branch_limit: int = 8
1950                                     ) -> float:
1951         """Compute expected value of taking 'action' in 'game' for this
1952             player.
1953
1954         Uses execute_spectrum when available then adds a rollout estimate
1955             for depth-1.
1956         """
1957         color = self._get_player_color()
1958
1959         # Quick boosts for robber/knight/dev before heavy sim
1960         name = self._safe_action_name(action)
1961         preboost = 0.0
1962         try:
1963             if any(tok in name for tok in ("move_robber", "robber")):
1964                 preboost += self.evaluate_robber_action(action, game,
1965                                                         color)
1966             if any(tok in name for tok in ("knight", "play_knight")):
1967                 preboost += self.evaluate_play_knight(action, game, color)
1968             if any(tok in name for tok in ("buy_dev", "buycard", "
1969                                         buy_dev_card")):
1970                 try:
1971                     dev_ev = self.dev_card_ev_estimate(game, color)
1972                     preboost += dev_ev * self.DEV_EV_SCALE
1973                 except Exception:
1974                     # fallback small preboost
1975                     preboost += 20.0
1976         except Exception:
1977             preboost += 0.0
1978
1979         branches = None
1980         try:
1981             branches = execute_spectrum(game, action)
1982             if not branches:
1983                 raise RuntimeError("execute_spectrum returned no branches
1984                                     ")
1985         except Exception as e_s:
1986             self.debug_print("FooPlayer: execute_spectrum failed or
1987                             unavailable for action;
1988                             trying deterministic.
1989                             Error:", e_s)
1990
1991             try:
1992                 branches = execute_deterministic(game, action)
1993                 if not branches:
1994                     raise RuntimeError("execute_deterministic returned no
1995                                         outcomes")
1996             except Exception as e_d:
1997                 self.debug_print("FooPlayer: Both execute_spectrum and
1998                                 execute_deterministic
1999                                 failed for action.
2000                                 Errors:", e_s, e_d)
2001
2002             return float("-inf")
2003
2004         # Limit branches to keep runtime bounded
2005         if len(branches) > per_action_branch_limit:
2006             branches = sorted(branches, key=lambda bp: float(bp[1]),
2007                               reverse=True)[:
2008                                   per_action_branch_limit]

```

```

1998     expected = 0.0
1999     total_prob = 0.0
2000     rollout_depth = max(0, self.ROLLOUT_DEPTH - 1)
2001     for (out_game, prob) in branches:
2002         try:
2003             # For buy_dev actions, if the branch encodes a known draw
2004                 outcome, we could
2005                 refine.
2006             # In absence of explicit draw info, rely on
2007                 dev_ev_estimate as a
2008                 conservative proxy.
2009             immediate = self._evaluate_game_state(out_game, color)
2010             rollout_est = self.rollout_value(out_game, color,
2011                 rollout_depth,
2012                 initial=True)
2013             branch_val = 0.6 * immediate + 0.4 * rollout_est
2014         except Exception as e:
2015             self.debug_print("FooPlayer: evaluation failed for branch
2016                 , using heuristic.
2017                 Error:", e)
2018             branch_val = self._heuristic_value(out_game, color)
2019             expected += float(prob) * float(branch_val)
2020             total_prob += float(prob)
2021
2022     if total_prob > 0:
2023         expected = expected / total_prob
2024
2025     expected += preboost
2026     return float(expected)
2027
2028 # ----- Main decision function -----
2029 def decide(self, game: Game, playable_actions: Iterable) -> Optional[
2030     object]:
2031     """Choose an action from playable_actions using phase-aware
2032         sampling + rollouts."""
2033     try:
2034         playable_actions = list(playable_actions)
2035         if not playable_actions:
2036             self.debug_print("FooPlayer: No playable actions
2037                 available, returning
2038                 None")
2039
2040             return None
2041
2042         color = self._get_player_color()
2043         phase = self.get_game_phase(game, color)
2044
2045         # Prefilter candidate actions
2046         candidates = self.prefilter_actions(playable_actions, game,
2047             color)
2048
2049         # Cap to MAX_SIMULATIONS
2050         if len(candidates) > self.MAX_SIMULATIONS:
2051             candidates = candidates[: self.MAX_SIMULATIONS]
2052
2053         if not candidates:
2054             candidates = random.sample(playable_actions,
2055                 min(len(playable_actions), self.MAX_SIMULATIONS))
2056
2057         # Distribute simulation budget adaptively
2058         per_action_budget = max(1, self.SIMULATION_BUDGET //
2059             max(1, len(candidates)))
2060
2061         best_score = float("-inf")
2062         best_actions: List[Any] = []
2063         scores_debug: List[Tuple[float, Any]] = []

```

```

2052
2053     for a in candidates:
2054         try:
2055             score =
2056                 self._evaluate_action_expectation(game, a,
2057                     per_action_branch_limit=per_action_budget)
2058         except Exception as e:
2059             self.debug_print("FooPlayer: Exception
2060                 during action evaluation, skipping action.
2061                 Error:", e)
2062             score = float("-inf")
2063
2064         scores_debug.append((score, a))
2065
2066         if score > best_score + self.TOLERANCE:
2067             best_score = score
2068             best_actions = [a]
2069         elif abs(score - best_score) <= self.TOLERANCE:
2070             best_actions.append(a)
2071
2072     # If no action had a finite score, fallback to first playable
2073     # action
2074     if not best_actions:
2075         self.debug_print("FooPlayer: All evaluations failed,
2076             defaulting to first
2077             playable action")
2078
2079     return playable_actions[0]
2080
2081     # Epsilon-greedy randomness to reduce predictability
2082     chosen: Any
2083     scores_debug.sort(key=lambda x: x[0], reverse=True)
2084     if random.random() < self.EPSILON_GREEDY and len(scores_debug
2085         ) >= 2:
2086         # pick from top-3 weighted by score (or fewer if not
2087         # available)
2088         top_k = scores_debug[: min(3, len(scores_debug))]
2089         weights = [max(0.0, s - top_k[-1][0] + 1e-6) for (s, a)
2090             in top_k]
2091         total_w = sum(weights)
2092         if total_w > 0:
2093             r = random.random() * total_w
2094             cum = 0.0
2095             for w, (_s, a) in zip(weights, top_k):
2096                 cum += w
2097                 if r <= cum:
2098                     chosen = a
2099                     break
2100             else:
2101                 chosen = top_k[0][1]
2102         else:
2103             chosen = scores_debug[0][1]
2104     if self.DEBUG:
2105         self.debug_print(f"FooPlayer: EPSILON pick
2106             triggered, chosen alternate action {chosen}")
2107     return chosen
2108
2109     # If tie, break ties preferring
2110     settlement/road/resource diversity improvements
2111     if len(best_actions) > 1:
2112         tie_metrics = []
2113         for a in best_actions:
2114             try:
2115                 metric = 0.0
2116                 metric += self.settlement_potential(a,

```

```

2106         game, color)
2107         metric +=
2108             self.road_connection_potential(a, game,
2109             color)
2110         # small production proxy via heuristic
2111         metric += 0.01 *
2112             self._heuristic_value(game, color)
2113         tie_metrics.append((metric, a))
2114     except Exception:
2115         tie_metrics.append((0.0, a))
2116     tie_metrics.sort(key=lambda x: x[0],
2117                     reverse=True)
2118     # pick the top metric actions (could still be
2119     multiple)
2120     top_metric = tie_metrics[0][0]
2121     filtered = [a for (m, a) in tie_metrics if abs(m
2122     - top_metric) <= self.TOLERANCE]
2123     if filtered:
2124         chosen = random.choice(filtered)
2125     else:
2126         chosen = random.choice(best_actions)
2127 else:
2128     chosen = best_actions[0]
2129
2130 # Debug logging: phase and top candidates
2131 if self.DEBUG:
2132     self.debug_print(f"FooPlayer: Phase={phase},
2133     SettlementsTarget=
2134     {self.TARGET_SETTLEMENTS_EARLY}")
2135     topn = scores_debug[:3]
2136     self.debug_print("FooPlayer: Top candidates:")
2137     for sc, act in topn:
2138         self.debug_print(f"    score={sc:.2f} action=
2139         {act}")
2140
2141     self.debug_print(f"FooPlayer: Chosen action {chosen}
2142     with expected score {best_score}")
2143     return chosen
2144 except Exception as e:
2145     # Protect against unexpected errors
2146     print("FooPlayer: Unexpected error in decide(),
2147     defaulting to first playable action. Error:", e)
2148     try:
2149         return list(playable_actions)[0]
2150     except Exception:
2151         return None

```

A.2 HEXMACHINA (NO DISCOVERY) SHALLOW STRATEGY

```

2150 import random
2151 import statistics
2152 from .adapters import (
2153     Game, Player, Color
2154 )
2155
2156 # Hyperparameters for this player. Tweak across evolutions.
2157 K_ROLLOUTS = 0 # rollouts disabled in this adapter-limited
2158                 implementation
2159 MAX_ROLLOUT_DEPTH = 10 # not used currently; kept for future use
2160 MAX_ACTIONS_TO_EVALUATE = 12
2161 DEBUG = True

```

```

2160
2161 class FooPlayer(Player):
2162     """A stronger FooPlayer that performs a 1-ply lookahead and evaluates
2163     the immediate successor state using a robust, defensive static
2164         evaluator.
2165
2166     Notes on integration with adapters.py:
2167     - We only use the thin adapter surface exported above (Game, Player,
2168         Color).
2169     - We call game.copy() to create hypothetical states and game.execute(
2170         action)
2171         to apply actions to those copies. We avoid calling any non-exported
2172         adapter helpers so this file remains compatible with the framework.
2173
2174     Limitations and rationale:
2175     - The adapters surface available in this environment does not
2176         explicitly
2177         expose helper functions for enumerating playable actions from an
2178         arbitrary game object (those are provided to decide() by the
2179         harness).
2180     Because of this we cannot reliably perform multi-step random
2181         rollouts
2182         (we cannot ask the engine for "playable_actions" inside the player
2183         for
2184         subsequent turns). Attempting to call hypothetical internal APIs
2185         would
2186         risk using non-portable / unsupported functions.
2187     - To still fix the key flaw (always pick the first action) we
2188         implement a
2189         1-ply lookahead over a sampled set of candidate actions and
2190         evaluate the
2191         successor state with a robust static value function that inspects
2192         the
2193         game.state. This is a significant upgrade over the previous
2194         behavior
2195         and provides a solid foundation for future rollout-based evolution.
2196     """
2197
2198 def __init__(self, name=None):
2199     super().__init__(Color.BLUE, name)
2200
2201 def decide(self, game, playable_actions):
2202     """Choose an action from playable_actions.
2203
2204     Strategy implemented:
2205     - If there are many playable actions, randomly sample up to
2206         MAX_ACTIONS_TO_EVALUATE actions to limit computation.
2207     - For each candidate action, copy the game, execute the action on
2208         the
2209         copy, and evaluate the resulting state with _evaluate_state().
2210     - Choose the action with the highest evaluation. Break ties
2211         randomly.
2212
2213     The evaluation is defensive: it attempts multiple common access
2214     patterns to extract victory points and common counts (settlements
2215         ,
2216     cities, roads). If extraction fails, the evaluator falls back to
2217         0.
2218
2219     Args:
2220         game (Game): complete game state. read-only. Use game.copy()
2221             to
2222             create hypothetical states.
2223         playable_actions (Iterable[Action]): legal options for this
2224             turn.

```

```

2214 Returns:
2215     action: chosen element of playable_actions, or None if no
2216             options.
2217 """
2218 # Defensive: if no actions available, return None
2219 if not playable_actions:
2220     if DEBUG:
2221         print('FooPlayer.decide: no playable_actions -> returning
2222               None')
2223     return None
2224
2225 # Convert playable_actions to a list so we can sample and index
2226 try:
2227     actions = list(playable_actions)
2228 except Exception:
2229     # If iterable cannot be converted, fall back to returning
2230     # first
2231     if DEBUG:
2232         print('FooPlayer.decide: playable_actions not list-like;
2233               defaulting to first')
2234     try:
2235         return playable_actions[0]
2236     except Exception:
2237         return None
2238
2239 # Sample candidate actions if there are too many
2240 if len(actions) > MAX_ACTIONS_TO_EVALUATE:
2241     candidates = random.sample(actions, MAX_ACTIONS_TO_EVALUATE)
2242     if DEBUG:
2243         print(f'FooPlayer.decide: sampled {len(candidates)} of {
2244               len(actions)}
2245               actions to evaluate')
2246 else:
2247     candidates = actions
2248     if DEBUG:
2249         print(f'FooPlayer.decide: evaluating all {len(candidates)}
2250               actions')
2251
2252 # Evaluate each candidate action by applying it to a copy of the
2253 # game
2254 scores = [] # list of (action, score)
2255 for i, action in enumerate(candidates):
2256     try:
2257         # Copy the game to avoid mutating the original
2258         new_game = game.copy()
2259
2260         # Apply the candidate action on the copied game.
2261         # The standard Game API exposes execute(action) to apply
2262         # an action.
2263         # We try both .execute and .apply for defensive
2264         # compatibility.
2265         executed = False
2266         try:
2267             new_game.execute(action)
2268             executed = True
2269         except Exception:
2270             # Some versions may expose a differently named method
2271             .
2272             try:
2273                 new_game.apply(action)
2274                 executed = True
2275             except Exception:
2276                 executed = False

```

```

2268
2269         if not executed:
2270             # If we couldn't apply the action on the copy, mark
2271                 it as
2272             # very poor and continue.
2273             if DEBUG:
2274                 print(f'FooPlayer.decide: failed to execute
2275                     candidate
2276                     action {i};
2277                     marking
2278                     score -inf')
2279             scores.append((action, float('-inf')))
2280             continue
2281
2282             # Evaluate the successor state
2283             score = self._evaluate_state(new_game)
2284             scores.append((action, score))
2285             if DEBUG:
2286                 print(f'FooPlayer.decide: action #{i} -> score {score
2287                     }')
2288
2289         except Exception as e:
2290             # Catch-all: do not let the player crash the harness.
2291                 Penalize
2292             # the action and continue evaluating others.
2293             if DEBUG:
2294                 print(f'FooPlayer.decide: exception while evaluating
2295                     action #{i}: {e}
2296                     ! Marking -inf')
2297             scores.append((action, float('-inf')))
2298
2299         # Choose the best action. If all are -inf or evaluation failed,
2300             fall back
2301         # to the original first-action policy.
2302         if not scores:
2303             if DEBUG:
2304                 print('FooPlayer.decide: no scores produced -> defaulting
2305                     to first action')
2306             return actions[0]
2307
2308         # Compute the maximum score
2309         max_score = max(score for (_, score) in scores)
2310         # Filter all actions that have the max score (handle ties)
2311         best_candidates = [a for (a, s) in scores if s == max_score]
2312
2313         if not best_candidates or max_score == float('-inf'):
2314             # All evaluations failed; fallback
2315             if DEBUG:
2316                 print('FooPlayer.decide: all evaluations failed ->
2317                     defaulting to first
2318                     action')
2319             return actions[0]
2320
2321         chosen = random.choice(best_candidates)
2322         if DEBUG:
2323             try:
2324                 # Try to pretty-print a small summary for debugging
2325                 print(f'FooPlayer.decide: selected action -> {repr(chosen
2326                     )} with score {
2327                     max_score}')
2328             except Exception:
2329                 print('FooPlayer.decide: selected an action (repr failed)
2330                     ')
2331
2332         return chosen

```

```

def _evaluate_state(self, game):
    """Static evaluation of a game state from this player's
        perspective.

    The evaluator attempts multiple common access patterns to extract
    victory points and simple progress indicators (settlements,
        cities,
    roads). The returned score is primarily the victory points (
        higher is
    better). Secondary counts are used as small tiebreakers.

    This function is defensive to avoid attribute errors across
        different
    engine versions.

    Returns:
        float: heuristic score for the state (larger is better)
    """
    color = self.color
    vp = None
    settlements = None
    cities = None
    roads = None

    # Try a number of plausible attribute access patterns. Use try/
    except
    # blocks liberally because different engine versions expose
    different
    # structures.
    try:
        players = game.state.players
    except Exception:
        players = None

    # Attempt to access player state by Color key
    player_state = None
    if players is not None:
        try:
            player_state = players[color]
        except Exception:
            # Maybe players is a list keyed by integer colors
            try:
                idx = int(color)
                player_state = players[idx]
            except Exception:
                player_state = None

    # Extract victory points with common attribute names
    if player_state is not None:
        for attr in ('victory_points', 'victoryPoints', 'vp', 'points
            '):
            try:
                val = getattr(player_state, attr)
                # If it's a callable (method), call it
                if callable(val):
                    val = val()
                vp = int(val)
                break
            except Exception:
                vp = None

    # Try dictionary-style if attributes failed
    if vp is None:
        try:

```

```

2376         if isinstance(player_state, dict):
2377             for key in ('victory_points', 'vp', 'points'):
2378                 if key in player_state:
2379                     vp = int(player_state[key])
2380                     break
2381     except Exception:
2382         vp = None
2383
2384     # Extract simple asset counts to break ties
2385     for attr in ('settlements', 'settle_count', 'settlement_count', 'settlements', 'settlements'):
2386         try:
2387             val = getattr(player_state, attr)
2388             if callable(val):
2389                 val = val()
2390             settlements = int(val)
2391             break
2392         except Exception:
2393             settlements = None
2394
2395     for attr in ('cities', 'city_count'):
2396         try:
2397             val = getattr(player_state, attr)
2398             if callable(val):
2399                 val = val()
2400             cities = int(val)
2401             break
2402         except Exception:
2403             cities = None
2404
2405     for attr in ('roads', 'road_count'):
2406         try:
2407             val = getattr(player_state, attr)
2408             if callable(val):
2409                 val = val()
2410             roads = int(val)
2411             break
2412         except Exception:
2413             roads = None
2414
2415     # Fallbacks if extraction failed: try to compute from visible board pieces
2416     # (e.g., lengths of lists). This is optional and best-effort.
2417     if vp is None and players is not None:
2418         try:
2419             # If player_state contains lists of pieces, inspect lengths
2420             if isinstance(player_state, dict):
2421                 # Look for settlement/city lists
2422                 s = None
2423                 for key in ('settlements', 'settle_list'):
2424                     if key in player_state and isinstance(player_state[key], (list, tuple)):
2425                         s = len(player_state[key])
2426                         break
2427                 if s is not None:
2428                     settlements = settlements or s
2429                 # We intentionally do not try to derive vp from the board in a brittle way; leave vp as None and fall back to 0.
2430             except Exception:
2431                 pass

```

```
2430     # Final fallback: if we couldn't determine vp, set to 0
2431     if vp is None:
2432         vp = 0
2433
2434     # Build a composite score. Main contributor is victory points.
2435         Add
2436     # small weighted bonuses for settlements/cities/roads if
2437         available.
2438
2439     score = float(vp)
2440     if settlements is not None:
2441         score += 0.01 * float(settlements)
2442     if cities is not None:
2443         score += 0.02 * float(cities)
2444     if roads is not None:
2445         score += 0.005 * float(roads)
2446
2447     return score
2448
2449
2450
2451
2452
2453
2454
2455
2456
2457
2458
2459
2460
2461
2462
2463
2464
2465
2466
2467
2468
2469
2470
2471
2472
2473
2474
2475
2476
2477
2478
2479
2480
2481
2482
2483
```