

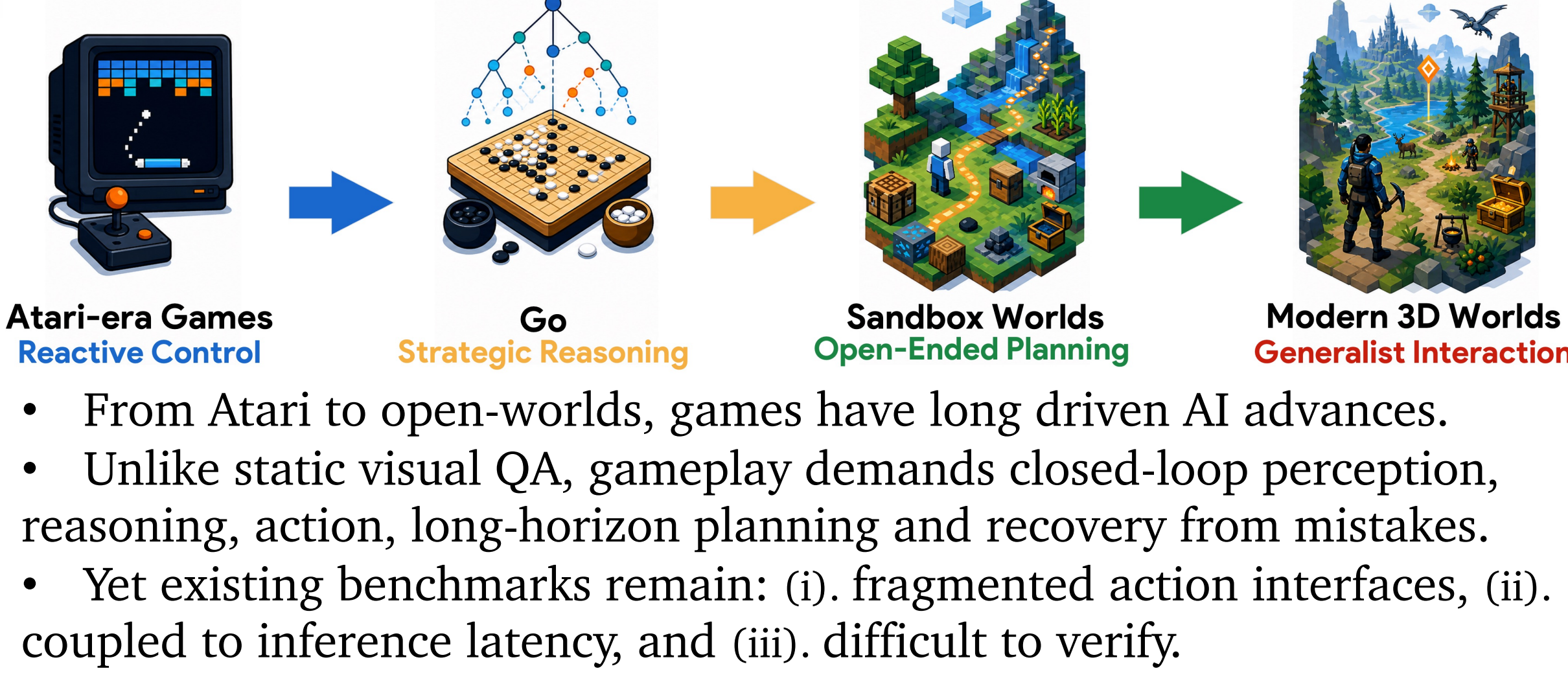


Towards Standardized and Verifiable Evaluation of Multimodal Game Agents

Mingyu Ouyang^{*1}, Siyuan Hu^{*1}, Kevin Qinghong Lin², Hwee Tou Ng^{†1}, Mike Zheng Shou^{†1}

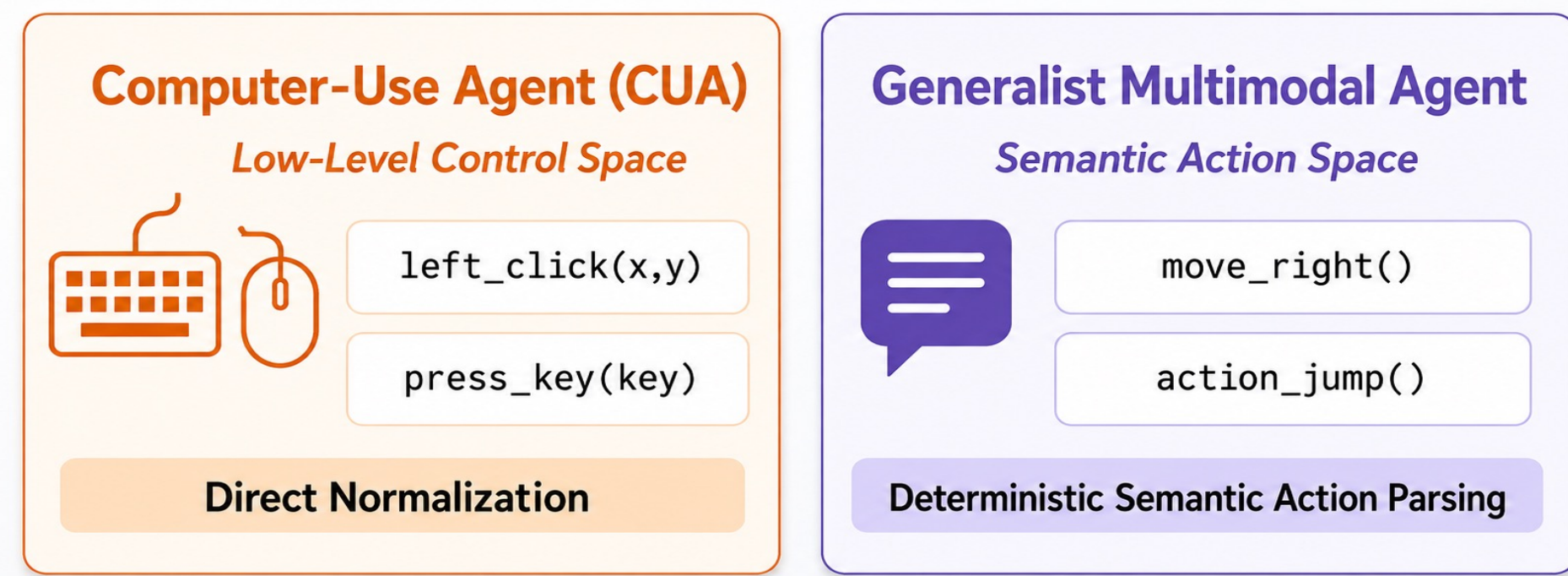
¹National University of Singapore, ²University of Oxford

Why Games?

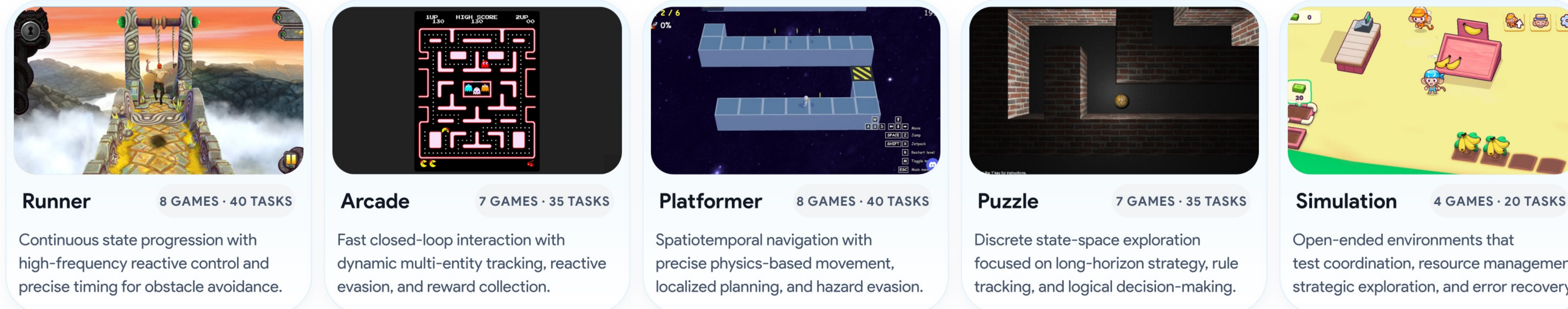


GameWorld: Standardized & Verifiable Evaluation

- Comparable across interfaces.** GameWorld standardizes both Computer-Use Agents and Generalist multimodal agents under shared browser environment.
- Environment diversity.** The suite spans 34 games across five genres to compare reactive control, spatial navigation, symbolic reasoning, and open-ended coordination under one protocol.
- Isolate decision quality and response speed.** The browser sandbox pauses game execution during model inference, so agents face identical dynamics regardless of latency.
- Verifiable outcome metrics.** Instead of visual heuristics or LLM-as-judge, GameWorld designs and reads serialized game state to compute progress and success directly from task-relevant variables.

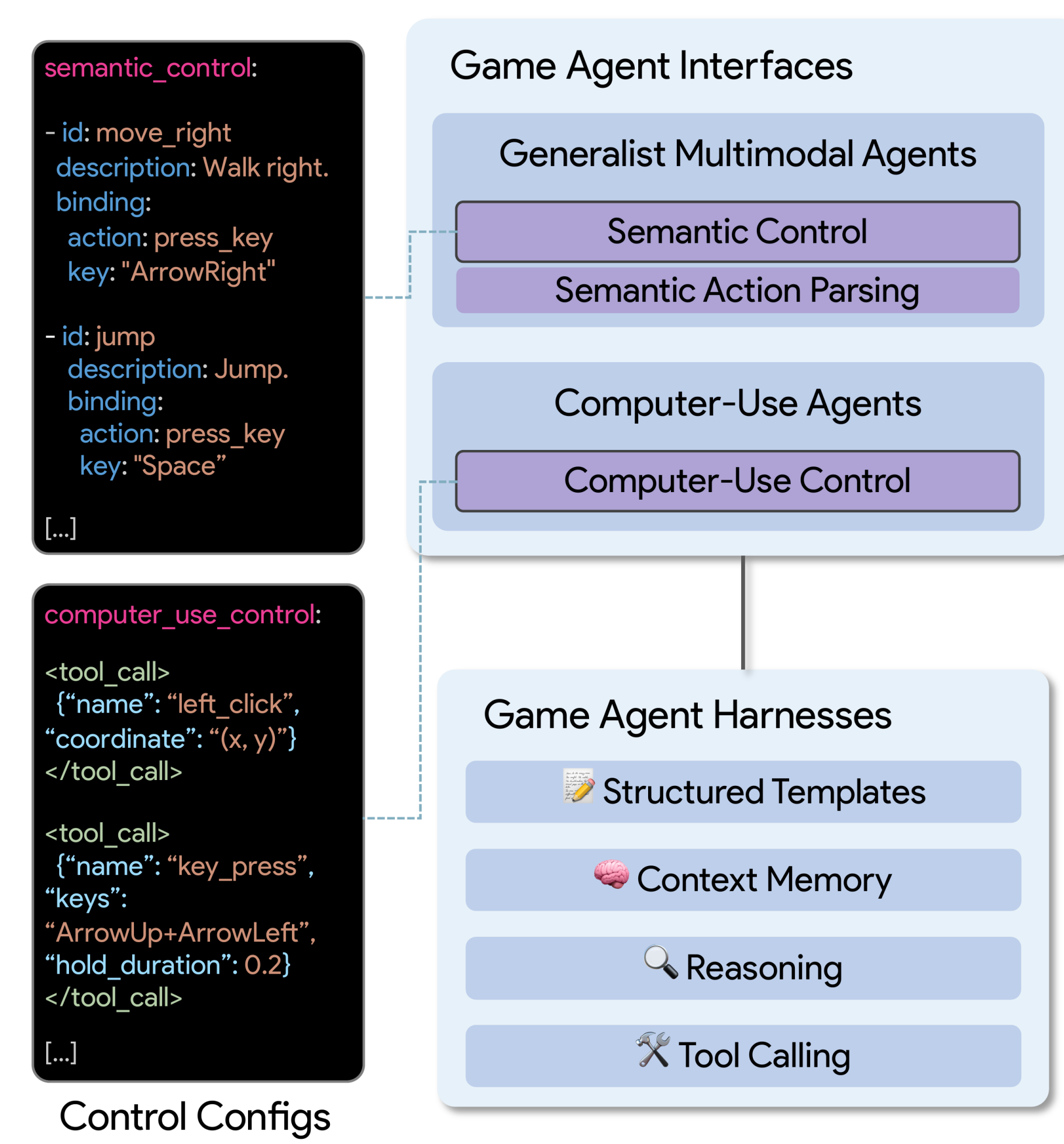


34 Games, 170 Verifiable Tasks

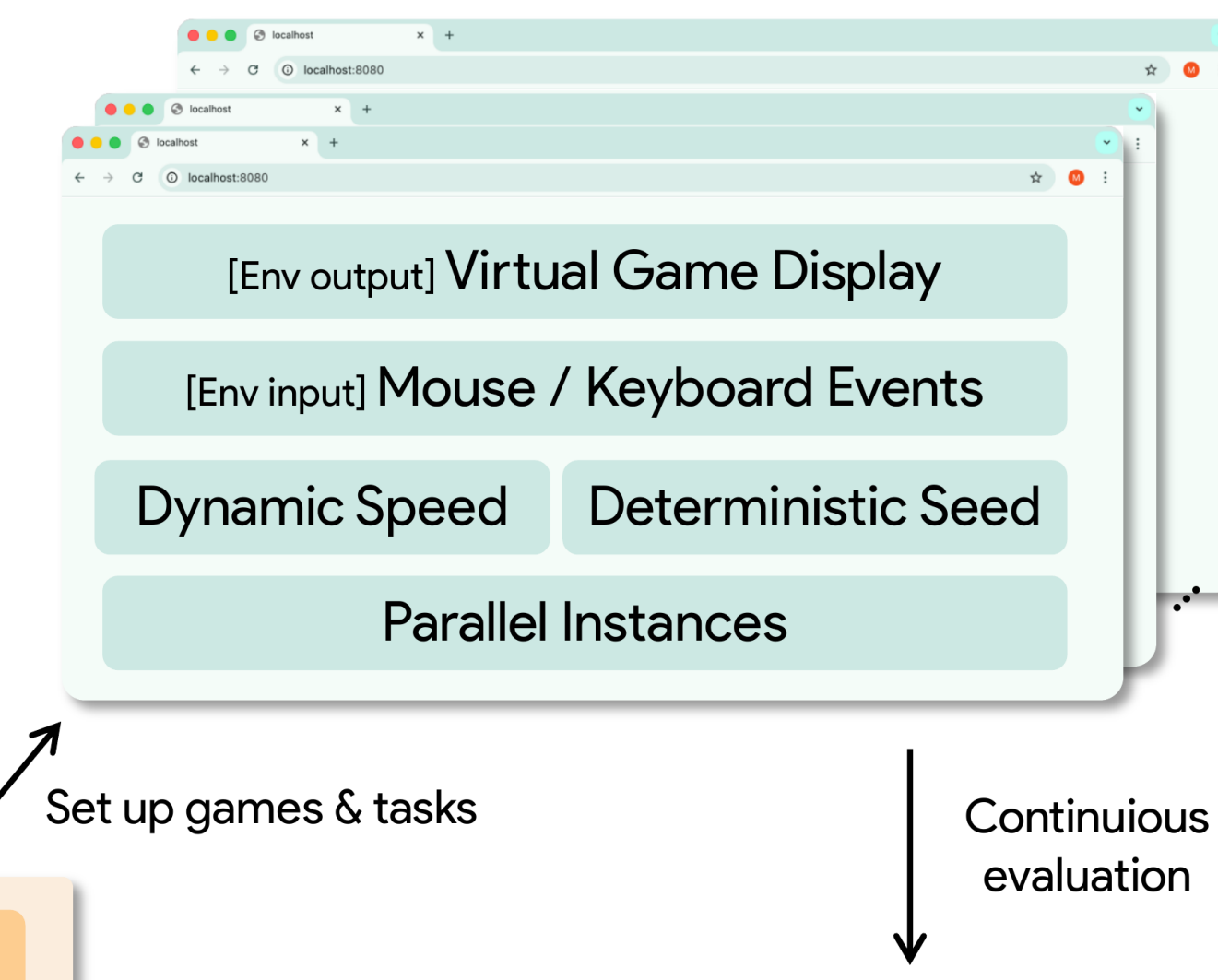


- Comprises 34 games and 170 verifiable tasks across five genres, designed to evaluate complementary gameplay capabilities.

(i) MLLMs as Game Agents



(ii) Browser-based Sandbox Environment



34 Game Roles and Rules
170 Task Instructions and Goals

(iii) Games & Tasks Library

```

role_config:
- game_role: "You are controlling Mario. Navigate through platforming levels by running, jumping, and avoiding enemies. [...]"
- game_rules: "You have 3 lives. Touching enemies from the side kills Mario. [...]"
task_config:
- task_instruction: "Collect 3 coins in this level."
evaluator_config:
- target_scoring_field: metrics.coins_collected
- target_score: 3
    
```

(iv) Outcome-based Evaluation

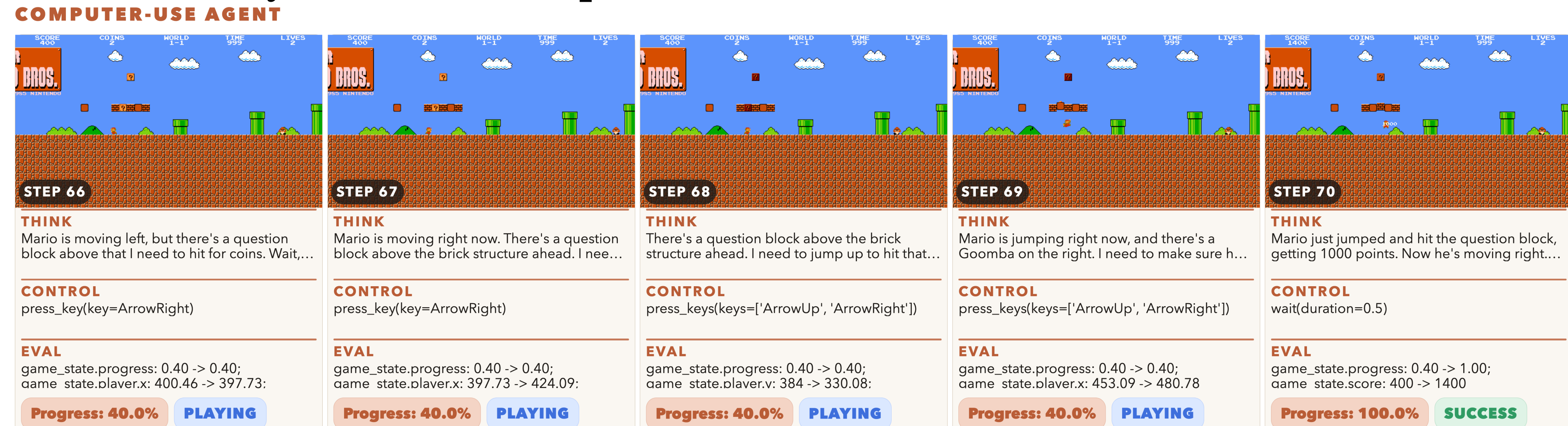
```

"game_state": {
  "score": 0, "level": "1-1", "progress": 0.35,
  "player": {"x": 98, "y": 384, "alive": true},
  "metrics": {"lives": 3, "coins_collected": 1, "distance": 33, "attempts": 1, "time_left_s": 999}, [...]
}
    
```

TASK OUTCOME:
PROGRESS: 0.33, SUCCESS: ✗

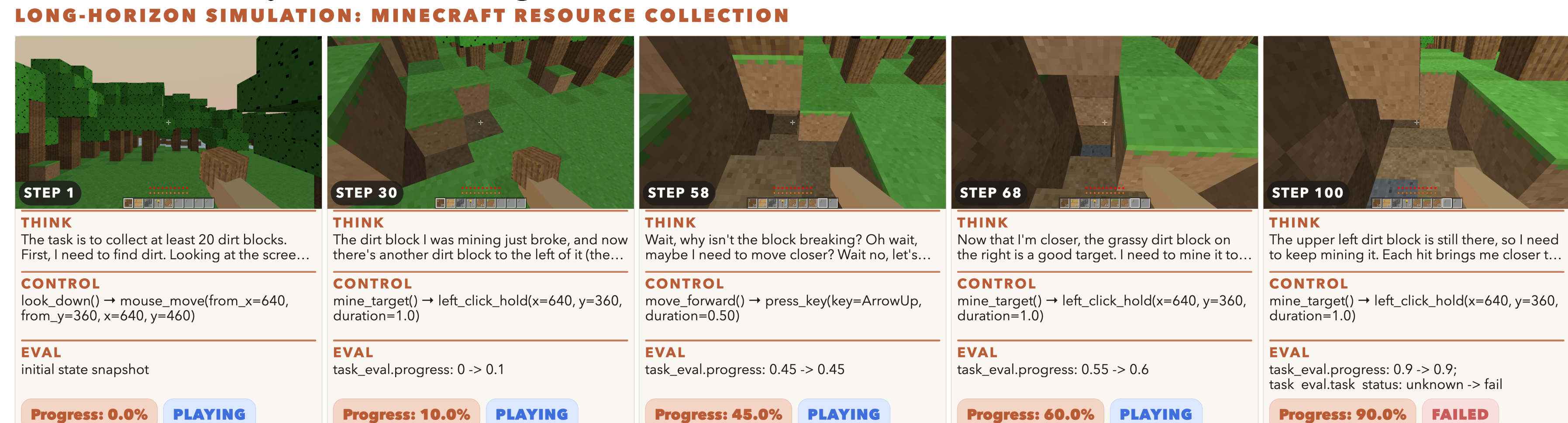
Framework Overview: GameWorld closes a controlled observation-action-verification loop for evaluating agents on games.

Case Study 1: Direct Computer-Use vs. Semantic Control Interface

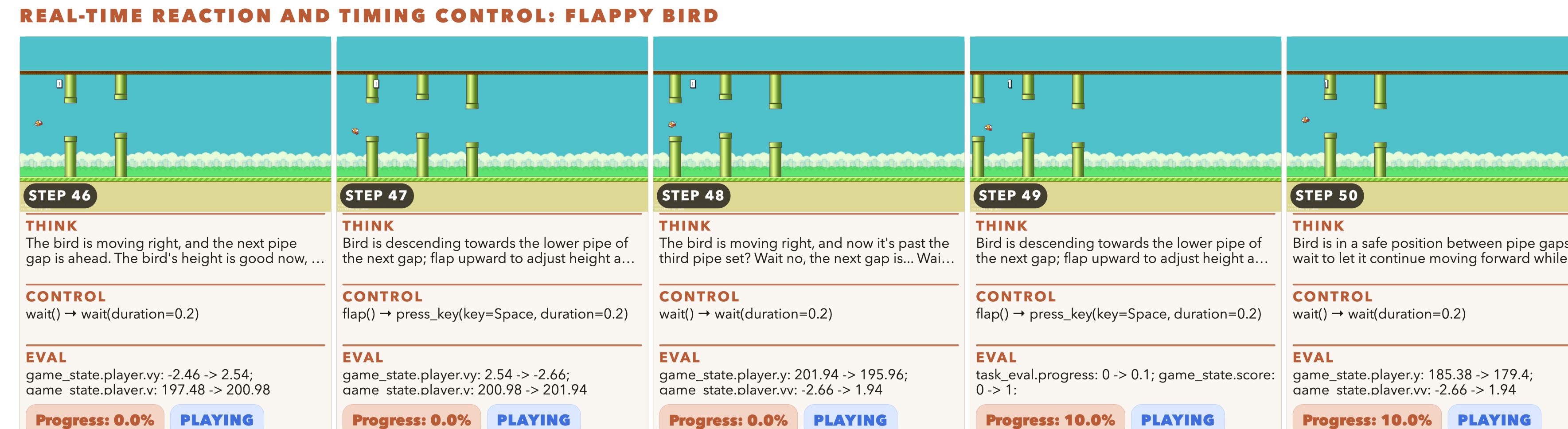


- Matched Mario trajectories: the **Computer-Use Agent** emits raw keyboard and mouse actions, while the **Generalist Agent** selects one of semantic actions that mapped deterministically to executable game controls.

Case Study 2 & 3: Long-Horizon Goal & Real-Time Reactive Control

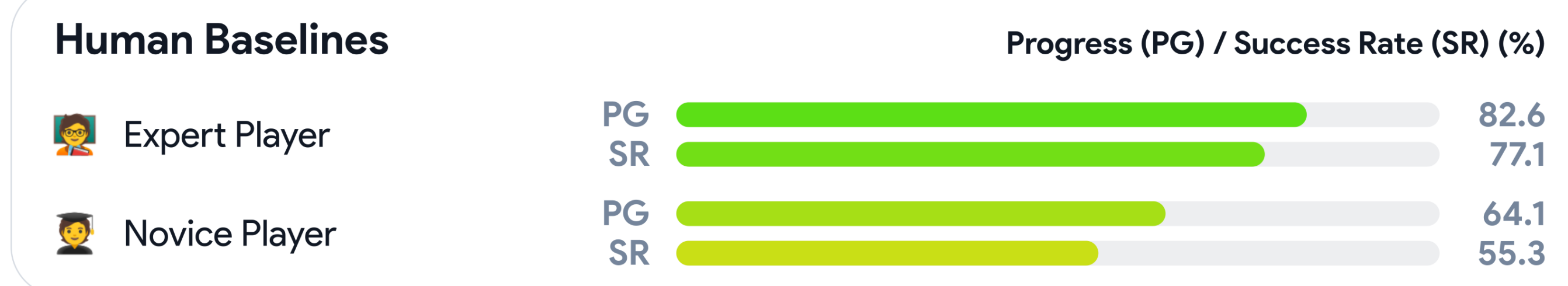
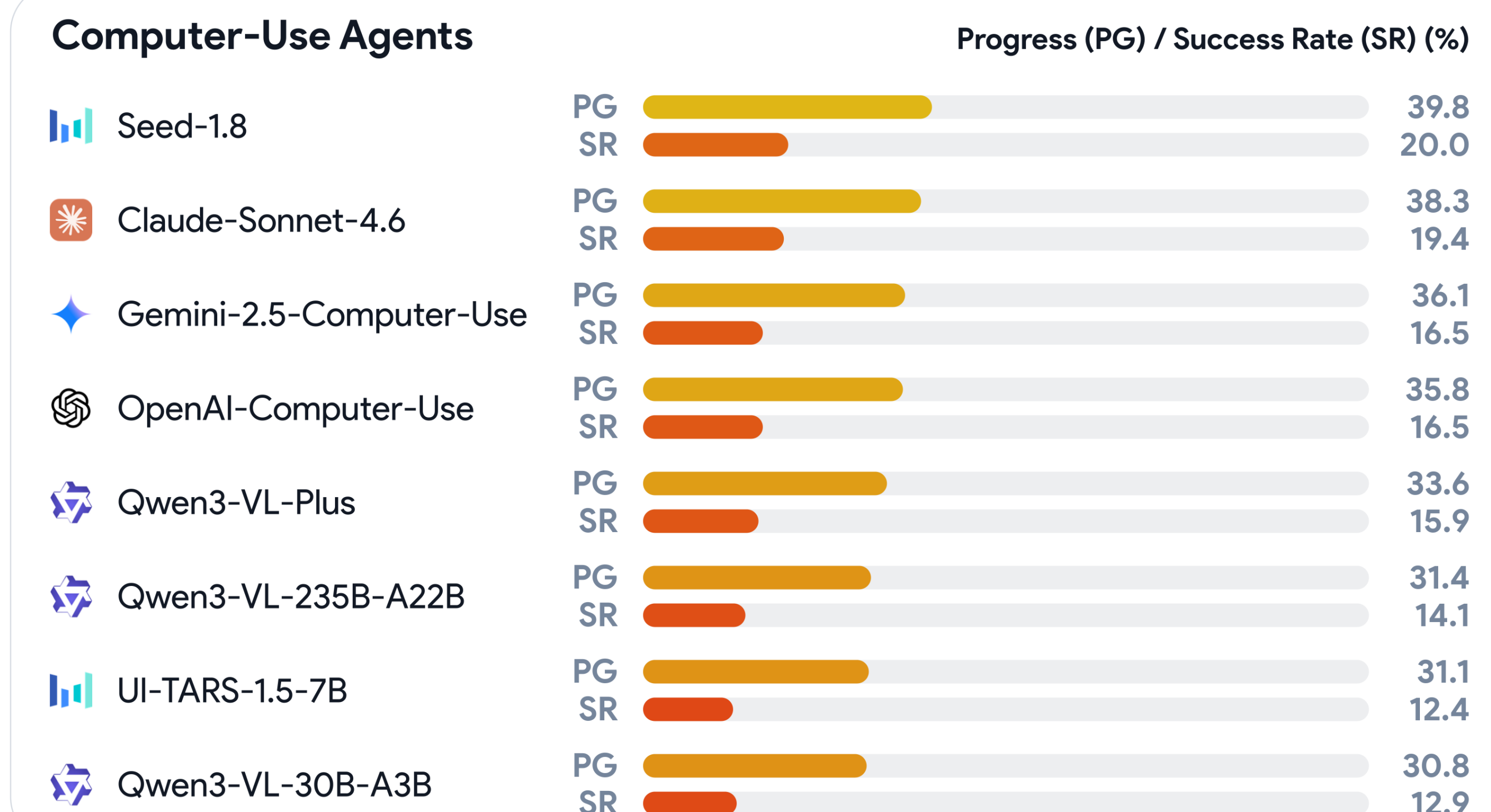
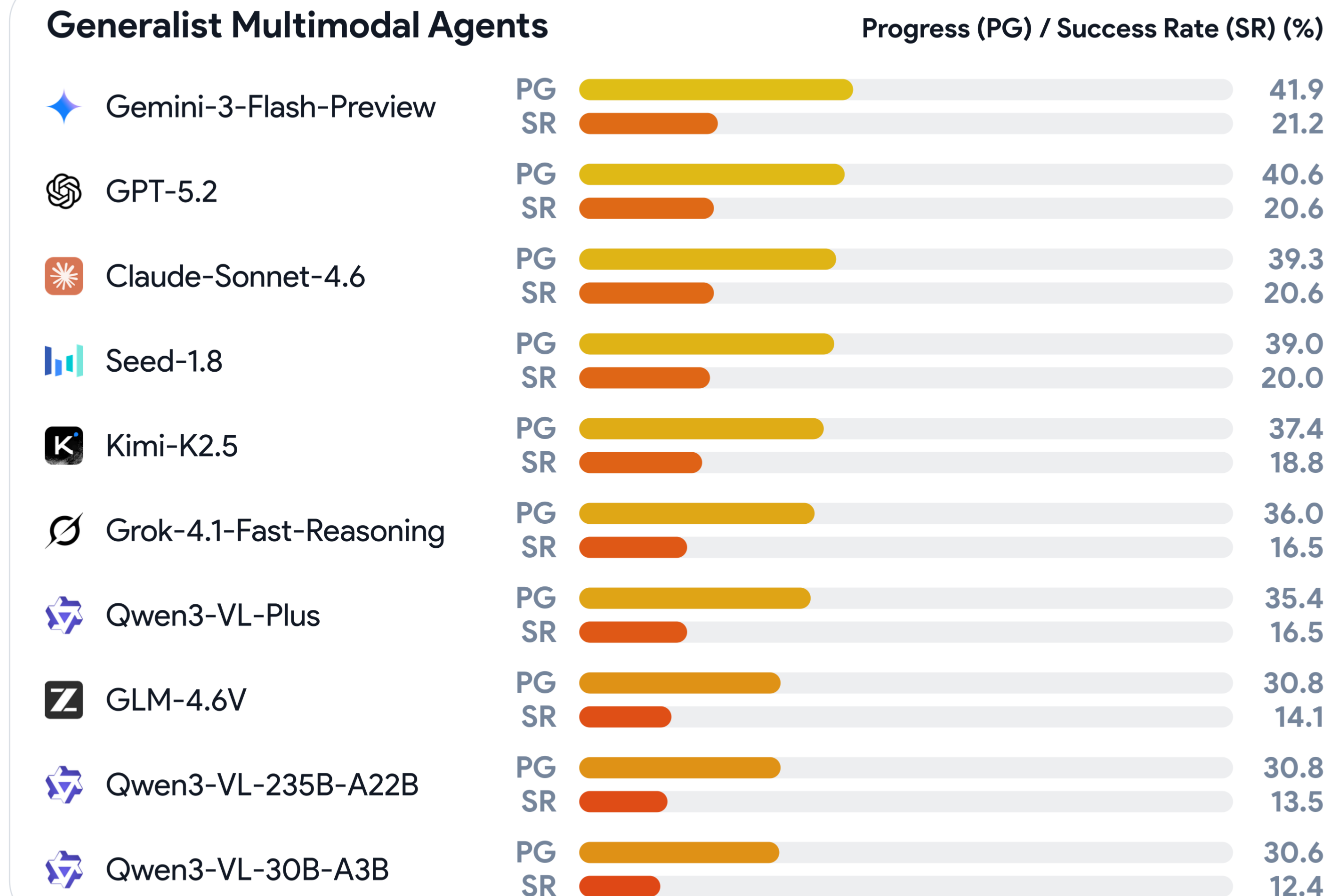


- Task: "Collect 20 dirt blocks." The agent reaches **90% progress** (18/20 collected) but misses task full completion.



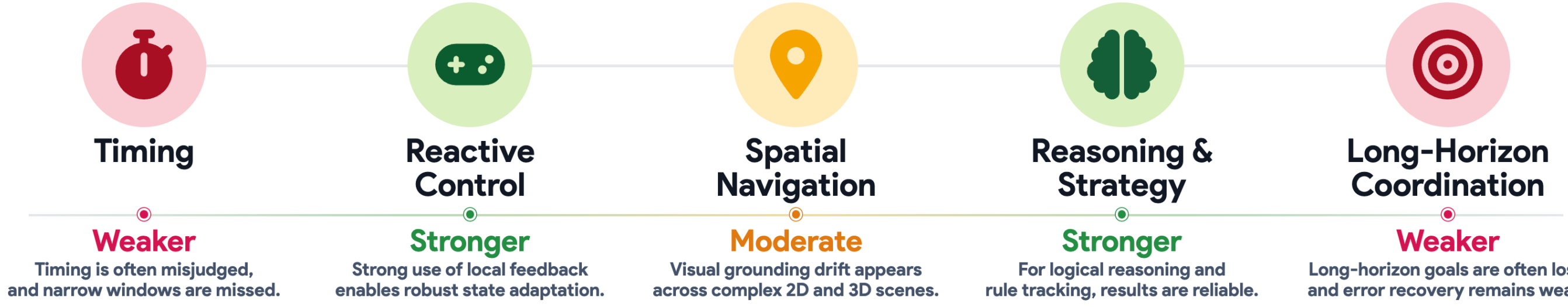
- Flappy Bird: visually consecutive frames require alternating wait/flap actions with precise timing.

Results & Analysis



- Best Generalist: **41.9% PG / 21.2% SR**; best CUA: **39.8% PG / 20.0% SR**.
- Both remain far below the novice human at **64.1% PG / 55.3% SR**.

Capability Bottlenecks



- Across games and agent interfaces, reactive control and symbolic reasoning are comparatively **stronger**; timing grounding, long-horizon completion, and coordination remain the **weaker** capabilities.

More detailed benchmark implementation and results:

