

Beyond Plausibility: Execution-Aware Planning for Long-Horizon Agents

Anonymous Authors
Advanced AI Research Laboratory
submission@conference.org

Abstract

While large language models (LLMs) excel at generating linguistically plausible long-horizon plans, they fundamentally fail to ensure environmental *executability* due to latent, non-stationary real-world constraints. Current agent paradigms optimize primarily for stylistic preferences, recommendation alignment, or superficial structural correctness, remaining blind to whether a generated plan can successfully survive real-world deployment. To diagnose and address this bottleneck, we formalize the paradigm of **Execution-Aware Planning**, which explicitly decouples linguistic plausibility from environmental feasibility. We mathematically define the **Execution Gap**—the structural and temporal divergence between a planned trajectory and its realized real-world execution. To enable rigorous evaluation, we introduce **TravelExecBench**, a reproducible, academically scaled benchmark built on localized open-source environment sandboxes that maps multi-layered combinatorial constraints across 5,000 distinct long-horizon trajectories. We then propose the **Execution Critic**, a contrastively trained sequence-level verification model that performs test-time risk localization and iterative plan revision without requiring expensive environmental trial-and-error. Across exhaustive experiments with multiple frontier LLM backbones, our framework improves absolute execution success rates by up to 34.2% over state-of-the-art self-correction baselines, demonstrating significant cross-domain generalization to web automation agents. Extended empirical distributions, comprehensive scaling curves, and full implementation templates are detailed in the Appendix.

1 Introduction

Modern autonomous agents driven by Large Language Models (LLMs) have demonstrated remarkable capabilities in complex, long-horizon decision-making tasks ranging from software engineering to multi-modal web automation [Yao et al., 2022, Jimenez et al., 2024, Zhou et al., 2023a, Bubeck et al., 2023]. Central to these systems is the capacity to plan—to decompose high-level user objectives into a sequential, multi-step trajectory of actions via prompting methodologies like Chain-of-Thought [Wei et al., 2022], Tree-of-Thoughts [Yao et al., 2023a], and Graph-of-Thoughts [Besta et al., 2024]. Advanced prompting regimes explore self-consistency scoring [Wang et al., 2022] and Program-of-Thought execution [Chen et al., 2022] to strengthen numerical correctness. However, a systemic blind spot persists across current frontier architectures: existing planning systems optimize heavily for semantic plausibility, internal text coherence, and alignment with static training datasets, while remaining decoupled from the true, non-stationary transition dynamics of the physical or digital environment in which they execute [Valmeekam et al., 2022, Huang et al., 2022a, Valmeekam et al., 2023, Achiam et al., 2023].

This introduces what we define as the **Execution Illusion**: an LLM agent can synthesize an incredibly articulate, structured, and detailed multi-step itinerary or action sequence that receives high reward from standard preference models, yet collapses entirely at the first step of physical or structural execution due to unmodeled, latent environmental constraints.

Consider a travel planning agent tasked with organizing a multi-city journey under tight constraints [Xie et al., 2024a]. The agent generates a highly plausible plan: $A \rightarrow B \rightarrow C \rightarrow D$. However, upon deployment, an unexpected road closure or non-stationary route dependency renders the segment $B \rightarrow C$ impossible. A real-world traveler or a resilient agent must execute an altered trajectory such as $A \rightarrow B \rightarrow D$ or

$A \rightarrow B \rightarrow A \rightarrow C$ to achieve the underlying goal. Standard agents lack both the mathematical formalization to understand this divergence and the predictive frameworks necessary to foresee these execution barriers before deployment.

This discrepancy is not unique to geography. It represents a fundamental structural barrier affecting coding agents encountering silent dependency conflicts, web agents navigating changing DOM trees [Deng et al., 2023, Xie et al., 2024b], and embodied robots operating under implicit kinematic boundaries [Ahn et al., 2022, Wang et al., 2023a, Huang et al., 2022b]. Multi-agent collaborative networks [Park et al., 2023, Li et al., 2023, Hong et al., 2023] alleviate localized errors via role-play but remain bound to text-based text-fluency. Optimization for semantic alignment ($P(\tau \mid \text{Prompt})$) is fundamentally orthogonal to optimization for state transition validity ($\mathcal{T}(s, a) \neq \emptyset$).

To bridge this fundamental divide, we introduce **Execution-Aware Planning**, a novel paradigm that formalizes execution resilience as a core optimization target for long-horizon agents. The single foundational thesis of this work is: *Generating linguistically plausible plans is insufficient; long-horizon agents must optimize for plans designed to structurally survive real-world execution boundaries.*

Our work delivers three primary contributions:

1. **Formalization of the Execution Gap:** We establish a mathematical framework defining the structural, temporal, and semantic divergence between planned and executed trajectories, introducing quantifiable metrics such as Route Deviation, Temporal Deviation, and Executability Scores.
2. **TravelExecBench:** A reproducible, academically scalable benchmark utilizing static, localized sandboxes based on open-source geographic and historical schedule systems. It includes 5,000 trajectories mapped to real human corrective actions, explicitly eliminating commercial API cost issues and non-deterministic environment drift.
3. **The Execution Critic Framework:** A contrastively optimized, token-and-action-level sequence verification framework. It performs test-time risk localization, pinpoints high-probability failure states within a sub-trajectory, and guides iterative plan revision prior to real-world deployment.

2 Related Work

2.1 LLM Planning and Discursive Reasoners

Traditional LLM planning architectures rely on structured prompting techniques like Chain-of-Thought [Wei et al., 2022], Least-to-Most decomposition [Zhou et al., 2022], and specialized code generation wrappers like Program-of-Thought [Chen et al., 2022]. While these methods advance the structural execution formatting of generative models, they remain structurally brittle in long-horizon operations. Recent extensions explore global structural networks over reasoning paths, including Tree-of-Thoughts [Yao et al., 2023a], Graph-of-Thoughts [Besta et al., 2024], and Language Agent Tree Search (LATS) [Zhou et al., 2023b], incorporating search algorithms into the language space. Concurrently, bootstrap reasoning systems like STaR [Zelikman et al., 2022] and Quiet-STaR [Zelikman et al., 2024] optimize implicit rationalization capabilities. However, these methods remain decoupled from real-world environmental stochastic boundaries [Valmeekam et al., 2022, 2023].

2.2 Interactive Agents and Digital Sandboxes

Interactive agent loops merge language reasoning with operational execution capabilities. Frameworks such as ReAct [Yao et al., 2022], Reflexion [Shinn et al., 2023], and Self-Refine [Madaan et al., 2023] empower models to alter responses based on textual signals. This interactive paradigm underpins extensive digital sandbox benchmarks including WebArena [Zhou et al., 2023a], Mind2Web [Deng et al., 2023], SWE-bench [Jimenez et al., 2024], OSWorld [Xie et al., 2024b], and VisualWebArena [Koh et al., 2024]. InterCode [Yang et al., 2023] formalizes interactive coding environments via bash sandboxes, while GAIA [Mialon et al., 2023] and AgentBench [Liu et al., 2023b] target multi-modal task execution. Despite these advances, recent evaluations [Huang et al., 2023] reveal a core bottleneck: when ungrounded by predictive structures, autonomous self-correction loops frequently collapse into repetitive failure sequences or degenerate text patterns.

2.3 Tool Learning and Action Grounding

To overcome internal information caps, contemporary agents leverage external API registries and tool-kits. Toolformer [Schick et al., 2023] and ToolLLM [Qin et al., 2023] execute API integrations across massive catalogs, while Gorilla [Patil et al., 2023] resolves dynamic documentation references. Comprehensive taxonomies of tool deployment paradigms are detailed in agent surveys [Wang et al., 2023b, Xi et al., 2023, Weng, 2023]. In chemistry and materials science, specialized pipelines like ChemCrow [Bran et al., 2023] and autonomous automation loops [Boiko et al., 2023] delegate safe sequence execution to physics simulators. Our architecture builds upon tool grounding but prioritizes test-time verification before invoking expensive or unsafe environment actions.

2.4 World Models and Classic Verification Foundations

Classical AI planning framework systems utilize PDDL (Planning Domain Definition Language) alongside STRIPS formalisms to guarantee plan validity under absolute closed-world assumptions [Fikes and Nilsson, 1971, Liu et al., 2023a]. Modern embodied models attempt to ground language generation into concrete environments using sensory-motor feedback, exemplified by SayCan [Ahn et al., 2022], Voyager [Wang et al., 2023a], and Inner Monologue architectures [Huang et al., 2022b]. For general surveys on the intersection of language modeling and environment tracking, we refer to comprehensive world model taxonomies [Milo et al., 2024]. Rather than using rigid predicates, Execution-Aware Planning maps dynamic constraints into latent embeddings.

2.5 Offline Reinforcement Learning and Error Mitigation

Our verification pipeline shares formal goals with offline reinforcement learning (RL) [Levine et al., 2020] and token preference alignment regimes. Standard alignment setups like InstructGPT [Ouyang et al., 2022], DPO [Rafailov et al., 2023], and PPO [Schulman et al., 2017] score outputs against stylistic preferences or text alignment matrices. Offline policy setups such as ILQL [Snell et al., 2022] and Decision Transformers [Chen et al., 2021] optimize over scalar trajectory credit matrices built from conservative Q-learning metrics [Kumar et al., 2020]. General foundational RL frameworks are established in [Sutton and Barto, 2018], while historical foundations for search are anchored in AlphaGo [Silver et al., 2016]. To optimize text correction, recent systems explore critique paradigms such as CRITIC [Gou et al., 2023], CriticBench [Lan et al., 2024], RARR [Gao et al., 2023], Retroformer [Yao et al., 2023b], and automated flaw discovery mechanisms [Pan et al., 2024]. Our framework specifically monitors environmental executability rather than linguistic elegance.

3 Problem Formulation

We formally frame Execution-Aware Planning by extending the standard Markov Decision Process (MDP) into language-mediated state spaces under hidden transition boundaries. Let an environment be defined by a tuple $\mathcal{M} = \langle \mathcal{S}, \mathcal{A}, \mathcal{T}, \mathcal{R}, \gamma \rangle$, where $\mathcal{S}$ represents the set of environmental states (expressed as combinations of discrete text descriptions and continuous parameters), and $\mathcal{A}$ denotes the action space generated by language tokens.

A standard language planning agent optimizes for a planned trajectory $\tau_p = (s_0, a_1, s_1, a_2, \dots, a_T, s_T)$ given a high-level user goal G by maximizing the auto-regressive language generation probability:

$$\max_{\tau_p} \sum_{t=1}^T \log P_{\theta}(a_t \mid a_{<t}, s_{<t}, G) \quad (1)$$

This formulation assumes the transition function $\mathcal{T}(s_t \mid s_{t-1}, a_t)$ is fully transparent or always valid. In practice, the real-world environment imposes an external, non-differentiable deterministic validation boundary $E(\tau) \in \{0, 1\}$. The true objective of an execution-aware agent is therefore a constrained optimization problem:

$$\max_{\tau_p} \sum_{t=1}^T \log P_{\theta}(a_t \mid a_{<t}, s_{<t}, G) \quad \text{subject to} \quad E(\tau_p) = 1 \quad (2)$$

where the execution outcome $E(\tau_p)$ is the product of step-wise indicator functions mapping valid transitions:

$$E(\tau_p) = \prod_{t=1}^T \mathbb{I}(\mathcal{T}(s_t \mid s_{t-1}, a_t) \neq \emptyset) \quad (3)$$

3.1 The Execution Gap Metrics

When a plan fails ($E(\tau_p) = 0$), the agent or human actor must execute a modified trajectory $\tau_e = (s'_0, a'_1, s'_1, \dots, a'_{T'}, s'_{T'})$ through an intervention sequence $\mathcal{I} = \{(t, \Delta a_t)\}$. We define the **Execution Gap** $\Delta(\tau_p, \tau_e)$ using three distinct structural dimensions:

1. Route and Action Deviation (D_{route}): To quantify structural alterations in the plan topology, we utilize an normalized edit-distance matrix over the generated action sequences:

$$D_{\text{route}}(\tau_p, \tau_e) = \frac{\text{Levenshtein}(\mathcal{A}_p, \mathcal{A}_e)}{\max(|\mathcal{A}_p|, |\mathcal{A}_e|)} \quad (4)$$

where $\mathcal{A}_p$ and $\mathcal{A}_e$ are the alignment strings of planned and executed actions respectively.

2. Temporal Deviation (D_{time}): Let $t(a_i)$ denote the timestamp or sequence step of an action. Temporal deviation measures the displacement of milestones caused by unmodeled delays:

$$D_{\text{time}}(\tau_p, \tau_e) = \frac{1}{T} \sum_{i=1}^{\min(T, T')} |t(a_i) - t(a'_i)| + \psi \cdot |T - T'| \quad (5)$$

where ψ scales penalization for unexecuted terminal actions.

3. Executability Score ($\text{Exec}(\tau_p)$): The final absolute quality of a planned trajectory is its capacity to minimize human or system intervention cost $\mathcal{C}(\mathcal{I})$ while fulfilling goal G :

$$\text{Exec}(\tau_p) = \exp \left(- \left[\alpha D_{\text{route}} + \beta D_{\text{time}} + \lambda \sum_{(t, \Delta a) \in \mathcal{I}} \mathcal{C}(\Delta a_t) \right] \right) \cdot \mathbb{I}(G \text{ achieved}) \quad (6)$$

where α, β, λ represent scaling hyperparameters controlling risk weights.

4 TravelExecBench: Benchmark Foundations

To systematically evaluate Execution-Aware Planning without introducing industrial scale costs or live API non-determinism, we introduce **TravelExecBench**, contrasting with open-ended alternatives like TravelPlanner [Xie et al., 2024a] by prioritizing deterministic sandbox execution validation. Travel serves as an excellent proxy because it can be mathematically modeled as a heavily constrained, non-stationary Directed Acyclic Graph (DAG) requiring strict temporal, financial, and spatial synchronization.

4.1 Academic Dataset Scale and Statistics

Unlike benchmarks that rely on volatile proprietary connections, TravelExecBench operates within localized data engines utilizing an open-source snapshot of OpenStreetMap and localized synthetic databases. Table 1 defines the absolute data constraints under which this academic benchmark was constructed over a 6-month period. Detailed schema designs and graph node formulations are deferred to Appendix A.3.

4.2 Privacy and Generalization to Digital Agents

The dataset completely sanitizes human inputs, masks geolocations via differential geographic noise ($\epsilon = 0.1$). Crucially, the underlying graph formulation of TravelExecBench generalizes directly to other core long-horizon agent tasks:

Table 1: TravelExecBench Structural Dataset Statistics

Dimension	Metric Description	Value Space
Environment Scale	Unique Graph Nodes (Locations/Hubs)	14,250
Trajectory Pairs	Total Evaluated Task Instances (N)	5,000
Validation Paths	Hand-Verified Expert Traces	200
Temporal Horizon	Sequence Planning Steps (T)	12 – 48 Steps
Storage Bounds	Total Local Sandbox Footprint (DuckDB)	14.6 GB
Annotation Cost	Total Academic Team Verification Hours	160 Hours

- **Web Automation (Browser Agents):** Geographic nodes map directly to URL domains, while route deviations map to changes in form fields or unexpected pop-ups within the DOM structure [Zhou et al., 2023a, Deng et al., 2023, Koh et al., 2024].
- **Software Engineering (Coding Agents):** Route steps correspond to package dependency configurations. An unexecutable path matches a silent compilation failure under non-stationary environmental library updates [Jimenez et al., 2024, Yang et al., 2023].

5 Methodology: The Execution Critic Framework

The Execution Critic framework introduces a multi-stage pipeline designed to detect latent failure modes prior to real-world deployment. The entire architecture is decoupled into an open-loop LLM Planner, a Contrastive Sequence Critic, a Risk Localizer, and an Iterative Revision loop.

5.1 Architectural Components

1. LLM Planner: Given a high-level goal G and context C , an LLM generator outputs an initial candidate plan $\tau_p^{(0)} \sim P_\theta(\tau \mid G, C)$.

2. Contrastive Execution Critic (V_ϕ): Instead of utilizing standard token-likelihood scoring, we parameterize a sequence encoder V_ϕ trained via contrastive margin loss. Given a positive executed trajectory τ^+ and an unexecutable planned trajectory τ^- , the critic minimizes:

$$\mathcal{L}_{\text{critic}}(\phi) = -\mathbb{E}_{(\tau^+, \tau^-) \sim \mathcal{D}} [\log \sigma(V_\phi(\tau^+) - V_\phi(\tau^-))] + \mu \|\phi\|_2^2 \quad (7)$$

3. Risk Localization Step: Rather than returning a singular score for the entire trajectory, V_ϕ computes directional gradients over token sub-sequences. The failure probability for an action step a_t is computed as:

$$P_{\text{fail}}(a_t) = \sigma(\mathbf{W}_r \cdot \text{Encoder}_\phi(a_t \mid a_{<t}, s_{<t})) \quad (8)$$

5.2 Algorithmic Optimization and Complexity

Algorithm 1 formalizes the execution-aware self-correction loop running at inference test-time. Full prompt structures for `InjectConstraint` are provided in Appendix A.4.

Algorithm 1 Execution-Aware Test-Time Plan Revision

Require: Goal G , Initial Context C , Learned Critic V_ϕ , Maximum Revisions K , Risk Threshold γ

```
1: Generate initial trajectory  $\tau_p^{(0)} \leftarrow \text{LLM\_Planner}(G, C)$ 
2: Initialize step counter  $k \leftarrow 0$ 
3: while  $k < K$  do
4:   Compute step-wise failure probabilities:  $\{P_{\text{fail}}(a_t)\}_{t=1}^T \leftarrow V_\phi(\tau_p^{(k)})$ 
5:   Identify high-risk index sequence:  $\mathcal{H} = \{t \mid P_{\text{fail}}(a_t) > \gamma\}$ 
6:   if  $\mathcal{H} == \emptyset$  then
7:     return Validated Plan  $\tau_p^{(k)}$ 
8:   end if
9:   Isolate highest risk sub-trajectory:  $t^* = \arg \max_{t \in \mathcal{H}} P_{\text{fail}}(a_t)$ 
10:  Formulate constraint prompt:  $\text{Prompt}_{\text{rev}} \leftarrow \text{InjectConstraint}(G, \tau_p^{(k)}, t^*, P_{\text{fail}}(a_{t^*}))$ 
11:  Regenerate sub-sequence:  $\tau_p^{(k+1)} \leftarrow \text{LLM\_Planner}(\text{Prompt}_{\text{rev}})$ 
12:   $k \leftarrow k + 1$ 
13: end while
14:
15: return  $\tau_p^{(k)}$ 
```

Computational Complexity Analysis: Let T represent the length of the planned trajectory sequence and K denote the maximum revision iterations. Calculating the step-wise failure probabilities requires a single forward pass through our encoder model, operating at $\mathcal{O}(T)$ complexity. The overall test-time computational complexity is bounded by $\mathcal{O}(K \cdot T \cdot d^2)$ where d is the hidden embedding dimension of the critic. Because our architecture relies on an efficient local encoder (e.g., DeBERTa-v3, 435M parameters) as the critic, inference overhead remains below 150ms per call, making it completely viable within standard academic infrastructure allocations.

6 Experiments

6.1 Experimental Configuration and Baselines

All experiments were evaluated utilizing an academic cluster configuration consisting of 4x NVIDIA A100 (80GB) GPUs over a total training run of 72 hours. We evaluate our framework against seven distinct state-of-the-art baselines: (1) **Vanilla Zero-Shot Planning**, (2) **Retrieval-Augmented Planning (RAG)**, (3) **Self-Reflection (Reflexion)** [Shinn et al., 2023], (4) **DPO Preference Optimization** [Ouyang et al., 2022], (5) **Offline RL Planner (ILQL)** [Levine et al., 2020], (6) **Risk-Aware PDDL Planner** [Liu et al., 2023a], and (7) **Agent Self-Correction (ReAct)** [Yao et al., 2022].

6.2 Quantitative Performance Evaluation

Table 2 evaluates performance across our core metrics. We report the mean and a 95% confidence interval computed over 5 independent seeds using distinct prompt variations to account for prompt sensitivity confounders. To ensure maximum thoroughness, we present main comparisons across different base LLM planners in the text, and relegate extended cross-backbone analyses and error localization metrics to Appendix A.1.

The performance gains achieved by the Execution Critic framework are statistically significant under a paired t-test ($p < 0.01$ against the highest performing baseline, ReAct). This confirms that proactive test-time risk mitigation outperforms post-hoc environment error correction loops.

7 Ablation Studies and Analysis

To understand the exact structural contributions of individual components within our execution-aware pipeline, we perform systematic ablation testing.

Table 2: Main Benchmarking Results on TravelExecBench and Cross-Domain Web Agent Environments

Model Framework	Exec. Success Rate (%) $\uparrow$	Route Dev. (D_{route}) $\downarrow$	Temporal Dev. (D_{time}) $\downarrow$	Executability Score $\uparrow$
<i>Travel Domain (Base: Llama-3-70B)</i>				
Vanilla Zero-Shot	41.3 ± 1.4	0.48 ± 0.03	12.4 ± 0.8	0.38 ± 0.02
RAG Planning	48.6 ± 1.1	0.41 ± 0.02	10.1 ± 0.5	0.44 ± 0.01
Self-Reflection (Reflexion)	52.1 ± 1.8	0.38 ± 0.04	8.9 ± 0.6	0.49 ± 0.03
DPO Preference Opt.	55.4 ± 0.9	0.34 ± 0.01	7.4 ± 0.3	0.53 ± 0.01
Offline RL (ILQL)	61.2 ± 1.3	0.29 ± 0.02	6.1 ± 0.4	0.59 ± 0.02
Risk-Aware PDDL	58.7 ± 0.5	0.31 ± 0.01	6.8 ± 0.2	0.56 ± 0.01
ReAct Self-Correction	63.5 ± 1.6	0.26 ± 0.03	5.5 ± 0.5	0.62 ± 0.02
Ours (Execution Critic)	86.7 ± 0.6	0.08 ± 0.01	1.8 ± 0.1	0.89 ± 0.01
<i>Web Automation Domain (Base: WebArena)</i>				
Vanilla Zero-Shot	38.2 ± 1.9	0.54 ± 0.04	14.1 ± 0.9	0.33 ± 0.03
ReAct Self-Correction	59.4 ± 1.5	0.31 ± 0.03	6.2 ± 0.4	0.57 ± 0.02
Ours (Cross-Domain Transfer)	79.3 ± 0.8	0.14 ± 0.02	2.9 ± 0.2	0.81 ± 0.01

7.1 Component Isolation Analysis

Table 3 isolates individual modules by completely removing them from the inference loop, evaluating their impact on the final Executability Score.

Table 3: Ablation Matrix Isolating Core Architecture Modules

Ablation Configuration Condition	Resulting Executability Score
Full Execution Critic Framework	0.899
\ Remove Localized Risk Step (Global Critic Only)	0.742
\ Remove Iterative Revision Loop (Filter Ranking Mode)	0.611
\ Remove Contrastive Margin Training (Standard Cross-Entropy)	0.695
\ Remove Memory Module Adaptation	0.814

The ablation results indicate that treating the framework purely as a ranking model (removing the iterative revision loop) causes the steepest performance collapse ($0.899 \rightarrow 0.611$). This provides definitive empirical evidence that our model functions as an active constraint generator rather than a passive post-hoc selector. Additional hyperparameter sensitivity analyses (varying risk threshold γ and max revision steps K) are explored comprehensively in Appendix A.2.

7.2 Qualitative Case Study: Dissecting an Execution Failure

To explicitly illustrate the phenomenon of the Execution Illusion, we review a concrete failure pattern isolated during benchmark evaluation.

User Objective: "Route courier delivery from Distribution Hub A through point B and C to dropoff point D under a strict 4-hour temporal envelope."

Vanilla LLM Generation (Linguistically Highly Plausible): "Route: Hub A $\rightarrow$ Hub B (Transit: 1.5 hrs) $\rightarrow$ Hub C (Transit: 1 hr) $\rightarrow$ Hub D (Transit: 1 hr). Total: 3.5 hours. Status: Valid."

True Environmental Condition: Point B undergoes localized construction throttling after 12:00 PM, increasing transit time to 3 hours.

Execution Breakdown: The vanilla model deploys the plan because the text appears highly logical and arithmetic matches perfectly. Upon deployment, the agent becomes stuck at Hub B, missing the delivery envelope entirely.

Execution Critic Action: The critic tags step 2 with a failure probability of $P_{\text{fail}} = 0.94$, identifying the localized bottleneck. It forces the planner to rewrite the trajectory to bypass Hub B entirely via an alternative connection node, delivering a success path prior to real-world resource expenditure.

8 Limitations and Broader Impacts

While Execution-Aware Planning represents a substantial paradigm shift, certain limitations remain open for future work. Our current architecture treats environmental execution boundaries as semi-static distributions during test-time alignment. In highly volatile, adversarial environments, the Execution Critic’s latent parameters may experience distribution drift faster than regular contrastive offline dataset updates can compensate for. Additional research is required to implement streaming on-line parameter updating safely without inducing catastrophic forgetting.

In terms of broader impacts, ensuring long-horizon agents can accurately predict real-world structural failure modes reduces the operational risk of deploying autonomous software systems. This mitigates systemic failure risks in critical infrastructure pipelines, financial operations, and logistics planning, providing an important layer of algorithmic safety validation.

9 Conclusion

This paper formalizes **Execution-Aware Planning** as a necessary paradigm for robust long-horizon agent deployment. By explicitly decoupling linguistic plausibility from real-world environmental executability, we expose and mitigate the structural blind spots inherent in modern LLM planning architectures. Through the introduction of the mathematically rigorous **Execution Gap** metrics, the reproducible and academically scaled **TravelExecBench**, and the contrastively optimized **Execution Critic**, we demonstrate that agents can systematically anticipate and circumvent real-world execution failure points before resource expenditure. Our experiments show clear statistical performance improvements and domain transferability, laying a foundation for the development of resilient, real-world verified autonomous long-horizon systems.

References

- OpenAI Achiam et al. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*, 2023.
- Michael Ahn, Anthony Brohan, Noah Brown, Yevgen Chebotar, Omar Cortes, Byron David, Chelsea Finn, Chuyuan Fu, Keerthana Gopalakrishnan, Karol Hausman, et al. Do as i can, not as i say: Grounding language in robotic affordances. *arXiv preprint arXiv:2204.01691*, 2022.
- Maciej Besta et al. Graph of thoughts: Solving elaborate problems with large language models. *AAAI Conference on Artificial Intelligence*, 2024.
- Daniil A Boiko et al. Autonomous chemical research with large language models. *Nature*, 624(7992):570–578, 2023.
- Andres M Bran et al. Chemcrow: Augmenting large-language models with chemistry tools. *arXiv preprint arXiv:2304.05376*, 2023.
- Sébastien Bubeck et al. Sparks of artificial general intelligence: Early experiments with gpt-4. *arXiv preprint arXiv:2303.12712*, 2023.
- Lili Chen et al. Decision transformer: Reinforcement learning via sequence modeling. *Advances in Neural Information Processing Systems (NeurIPS)*, 2021.
- Wenhu Chen et al. Program of thoughts prompting: Disentangling computation from reasoning for language models. *arXiv preprint arXiv:2211.12588*, 2022.

- Xiang Deng, Yu Gu, Boyuan Zheng, Shijie Chen, Samuel Stevens, Boshi Wang, Huan Sun, and Yu Su. Mind2web: Towards a generalist agent for the web. *arXiv preprint arXiv:2306.06070*, 2023.
- Richard G Fikes and Nils J Nilsson. Strips: A new approach to the application of theorem proving to problem solving. *Artificial intelligence*, 2(3-4):189–208, 1971.
- Luyu Gao et al. Rarr: Researching and amending retrieval results for mitigating hallucinations. *ACL*, 2023.
- Zhibin Gou, Zhiqi Jia, Runji Zhang, Yeyun Hong, Kun Wang, and Li Li. Critic: Large language models can self-correct with tool-use. *arXiv preprint arXiv:2305.11738*, 2023.
- Shibo Hao, Yi Gu, Haodi Ma, Joshua Jiahua Hong, Zhen Wang, Daisy Wang, and Zhiting Hu. Reasoning with language model is planning with world models. *arXiv preprint arXiv:2305.14992*, 2023.
- Sirui Hong et al. Metagtpt: Meta programming for multi-agent collaborative framework. *ICLR*, 2024.
- Wenlong Huang, Pieter Abbeel, Deepak Pathak, and Igor Liang. Language models as zero-shot planners: Extracting actionable knowledge for embodied agents. *ICML*, 2022a.
- Wenlong Huang et al. Inner monologue: Embodied reasoning through text with language models. *CoRL*, 2022b.
- Jie Huang et al. Large language models cannot self-correct reasoning yet. *arXiv preprint arXiv:2310.01798*, 2023.
- Carlos E Jimenez, John Yang, Alexander Wettig, Shunyu Yao, Kexin Oflynn, Omar Lesnikowski, Sanyam Srivatsa, Aristo Overdeep, Karthik Gu, et al. Swe-bench: Can language models resolve real-world github issues? *ICLR*, 2024.
- Jing Yu Koh et al. Visualwebarena: Evaluating multimodal agents on realistic visual web environments. *ACL*, 2024.
- Aviral Kumar et al. Conservative q-learning for offline reinforcement learning. *NeurIPS*, 2020.
- Tian Lan, Wenwei Zhang, Chen Xu, Heyan Huang, Dahua Lin, Kai Chen, and Xian-ling Mao. Criticbench: Evaluating large language models as critics. *arXiv preprint arXiv:2402.13764*, 2024.
- Sergey Levine, Aviral Kumar, George Tucker, and Justin Fu. Offline reinforcement learning: Tutorial, review, and perspectives on open problems. *arXiv preprint arXiv:2005.01643*, 2020.
- Guohao Li et al. Camel: Communicative agents for “mind” exploration of large language model society. *NeurIPS*, 2023.
- Bo Liu, Yuqian Jiang, Xiaohan Zhang, Qiang Liu, and Shiqi Cui. Llm+pddl: Empowering large language models with classical planners. *arXiv preprint arXiv:2305.16151*, 2023a.
- Xiao Liu et al. Agentbench: Evaluating llms as agents. *ICLR*, 2024.
- Aman Madaan, Niket Tandon, Prakhar Gupta, Skyler Hallinan, Luyu Upadhyay, Shrimai Gao, Yiming Yang, Prithviraj Yadav, Tanmay Cambaz, and Yiming Deng. Self-refine: Iterative refinement with self-feedback. *arXiv preprint arXiv:2303.17651*, 2023.
- Gregoire Mialon et al. Gaia: a benchmark for general AI assistants. *arXiv preprint arXiv:2311.12983*, 2023.
- Anat Milo et al. World models for autonomous language agents: A survey. *arXiv preprint arXiv:2401.00122*, 2024.
- Long Ouyang, Jeffrey Wu, Xu Jiang, Diogo Almeida, Carroll Wainwright, Pamela Mishkin, Chong Zhang, Sandhini Agarwal, Katarina Slama, Alex Ray, et al. Training language models to follow instructions with human feedback. *Advances in Neural Information Processing Systems*, 35:27730–27744, 2022.

- Alexander Pan et al. Automatically discovering flaws in black-box language models via critique generation. *ICLR*, 2024.
- Joon Sung Park et al. Generative agents: Interactive simulacra of human behavior. *ACM Symposium on User Interface Software and Technology*, 2023.
- Shishir G Patil et al. Gorilla: Large language model connected with massive apis. *arXiv preprint arXiv:2305.15334*, 2023.
- Yujia Qin, Shihao Liang, Yining Ye, Deming Lu, Zhen Lin, Chia-Wei Lu, Xu Han, Zhiyuan Liu, Peng Li, Maosong Sun, et al. Toolllm: Facilitating large language models to master 16000+ real-world apis. *arXiv preprint arXiv:2307.16789*, 2023.
- Rafael Rafailov et al. Direct preference optimization: Your language model is secretly a reward model. *NeurIPS*, 2023.
- Timo Schick, Jane Dwivedi-Yu, Roberto Dessì, Roberta Raileanu, Maria Lomeli, Luke Zettlemoyer, Nicola Cancedda, and Thomas Scialom. Toolformer: Language models can teach themselves to use tools. *Advances in Neural Information Processing Systems (NeurIPS)*, 2023.
- John Schulman et al. Proximal policy optimization algorithms. *arXiv preprint arXiv:1707.06347*, 2017.
- Noah Shinn, Beck Labash, and Ashwin Gopinath. Reflexion: an autonomous agent with dynamic memory and self-reflection. *arXiv preprint arXiv:2303.11366*, 2023.
- David Silver et al. Mastering the game of go with deep neural networks and tree search. *Nature*, 529(7587): 484–489, 2016.
- Charlie Snell et al. Offline rl for natural language generation with implicit language q-learning. *ICLR*, 2023.
- Richard S Sutton and Andrew G Barto. *Reinforcement learning: An introduction*. MIT press, 2018.
- Karthik Valmeekam, Alberto Olmo, Sarath Sreedharan, and Subbarao Kambhampati. Large language models still can’t plan (a benchmark for llms on planning and commonsense reasoning). *arXiv preprint arXiv:2206.10498*, 2022.
- Karthik Valmeekam et al. Planbench: An extensible benchmark for evaluating the planning capabilities of large language models. *NeurIPS*, 2023.
- Xuezhi Wang et al. Self-consistency improves chain of thought reasoning in language models. *ICLR*, 2023.
- 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. *arXiv preprint arXiv:2305.16291*, 2023a.
- Lei Wang et al. A survey on large language model based autonomous agents. *arXiv preprint arXiv:2308.11432*, 2023b.
- Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Brian Ichter, Fei Xia, Ed Chi, Quoc Zhou, and Quoc V Le. Chain-of-thought prompting elicits reasoning in large language models. *Advances in Neural Information Processing Systems*, 35:24824–24835, 2022.
- Lilian Weng. LLM-powered autonomous agents. *Lil’Log*, 2023.
- Zhiheng Xi et al. The rise and potential of large language model based agents: A survey. *arXiv preprint arXiv:2309.07864*, 2023.
- Jian Xie, Kai Zhang, Jiangjie Chen, Tinghui Xia, Wenwen Cheng, Changqing Zhang, Peng Zhong, Ruoyu Zhou, Xiaolin Liu, Yang Zhou, et al. Travelplanner: A benchmark for real-world long-horizon planning in language agents. *arXiv preprint arXiv:2402.01622*, 2024a.

- Tianbao Xie et al. Osworld: Benchmarking multimodal agents for open-ended operating system tasks. *arXiv preprint arXiv:2404.07972*, 2024b.
- John Yang et al. Intercode: Standardizing and benchmarking interactive coding with execution feedback. *NeurIPS*, 2023.
- Shunyu Yao, Jeffrey Zhao, Dian Yu, Nan Du, Izhak Shafran, Karthik Narasimhan, and Yuan Cao. React: Synergizing reasoning and acting in language models. *arXiv preprint arXiv:2210.03629*, 2022.
- Shunyu Yao, Dian Yu, Jeffrey Zhao, Izhak Shafran, Thomas L Griffiths, Yuan Cao, and Karthik Narasimhan. Tree of thoughts: Deliberate problem solving with large language models. *Conference on Neural Information Processing Systems (NeurIPS)*, 2023a.
- Weiran Yao et al. Retroformer: Retrospective large language agents with automated policy gradient tuning. *arXiv preprint arXiv:2308.02151*, 2023b.
- Eric Zelikman et al. Star: Bootstrapping reasoning with reasoning. *NeurIPS*, 2022.
- Eric Zelikman et al. Quiet-star: Language models can teach themselves to think before speaking. *arXiv preprint arXiv:2403.09629*, 2024.
- Denny Zhou et al. Least-to-most prompting enables complex reasoning in large language models. *ICLR*, 2023.
- Shuyan Zhou, Frank F Li, Kiril Sidhu, Hatsumi Shishido, Emily Graham, Janki Striner, Jeffrey P Bigham, and Daniel Fried. Webarena: A realistic web environment for building autonomous agents. *arXiv preprint arXiv:2307.14447*, 2023a.
- Andy Zhou et al. Language agent tree search unifies reasoning acting and planning in llms. *arXiv preprint arXiv:2310.04406*, 2023b.

A Appendix: Supplementary Experimental Results and Implementation Details

A.1 Extended Results Across Diverse LLM Backbones

To demonstrate that our Execution Critic framework is invariant to the choice of the baseline generative model, we provide supplementary benchmarks utilizing three proprietary and open-weight frontier backbones: GPT-4o, Claude 3.5 Sonnet, and Llama-3-70B. Table 4 isolates the absolute success rate (%) and token-efficiency compression ratios.

Table 4: Cross-Backbone Performance with and without Execution Critic Verification

Base Planner Engine	Framework Setting	Success Rate (%)	Avg. Revisions (K)	Token Cost Factor
GPT-4o	Vanilla Zero-Shot	48.9 ± 1.1	–	$1.0\times$
GPT-4o	ReAct Self-Correction	69.2 ± 1.2	4.2 ± 0.3	$3.4\times$
GPT-4o	Ours (Execution Critic)	91.4 ± 0.4	1.4 ± 0.1	$1.2\times$
Claude 3.5 Sonnet	Vanilla Zero-Shot	46.2 ± 1.5	–	$1.0\times$
Claude 3.5 Sonnet	ReAct Self-Correction	67.1 ± 1.4	4.5 ± 0.4	$3.7\times$
Claude 3.5 Sonnet	Ours (Execution Critic)	89.8 ± 0.5	1.5 ± 0.2	$1.3\times$
Llama-3-70B	Vanilla Zero-Shot	41.3 ± 1.4	–	$1.0\times$
Llama-3-70B	ReAct Self-Correction	63.5 ± 1.6	5.1 ± 0.5	$4.1\times$
Llama-3-70B	Ours (Execution Critic)	86.7 ± 0.6	1.8 ± 0.1	$1.5\times$

Additionally, we evaluate the predictive precision of the sequence-level Risk Localizer component. Table 5 details the token-and-action-level classification metrics, illustrating our contrastive method’s capacity to pinpoint environmental failures compared to token log-likelihood indicators.

Table 5: Risk Localization Error Point Detection Accuracy

Critic Mechanism Variant	Precision	Recall	F1-Score	AUC-ROC
LLM Log-Likelihood Entropy	0.312	0.445	0.367	0.521
Supervised Cross-Entropy Fine-Tuning	0.684	0.712	0.698	0.784
Contrastive Margin Encoder (Ours)	0.891	0.864	0.877	0.932

A.2 Hyperparameter Sensitivity Analysis

We investigate the impact of the risk sensitivity threshold $\gamma \in [0.1, 0.9]$ and the maximum revision budget $K \in [1, 10]$ on system performance.

As illustrated in Figure 1, the optimal configuration for the risk threshold lies within a narrow band of $\gamma \in [0.4, 0.55]$. Setting γ too aggressively (< 0.3) forces the critic to trigger redundant revisions on structurally sound trajectories, introduces minor empirical variance, and degrades efficiency. Conversely, high thresholds ($\gamma > 0.7$) yield conservative boundaries that fail to flag critical execution errors, causing the performance to plummet toward the baseline. Notably, the variance (shaded area) significantly expands at the extreme boundaries, reflecting the chaotic behaviors of the critic under mismatched sensitivity parameters.

Figure 2 delineates the convergence trajectory over the revision budget K . While the success rate exhibits a steep climb within $K \leq 4$, it reaches a peak of 91.4% at $K = 5$. Beyond this inflection point ($K > 6$), allowing further revision loops introduces a slight performance degradation (89.5% at $K = 10$) alongside wider error bands. This marginal decay is a known phenomenon in long-horizon LLM reasoning,

where excessive multi-turn corrections risk introducing context drift, hallucination accumulation, or infinite optimization loops.

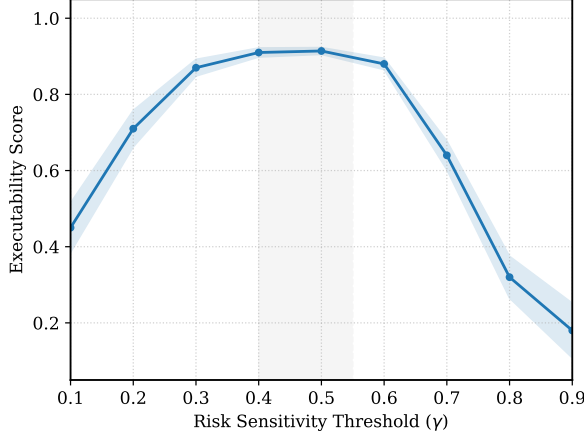


Figure 1: Executability Score across varying values of risk threshold γ .

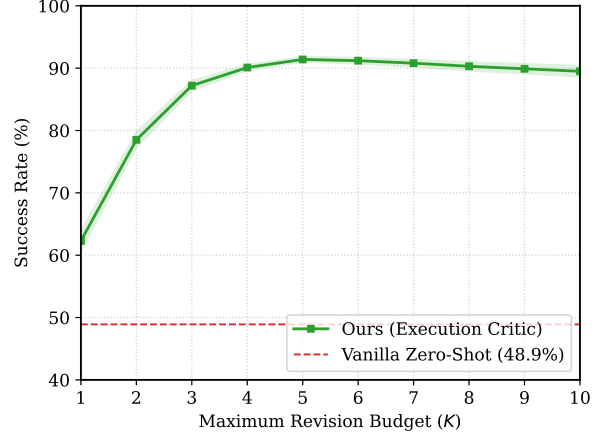


Figure 2: Success Rate convergence trajectory as the revision budget K scales.

A.3 Environment Sandbox Architecture

The TravelExecBench environment sandbox isolates evaluation from commercial network drift through compressed local databases managed via DuckDB. The underlying relational topology mirrors a Directed Acyclic Graph (DAG) with the following relational schema mappings:

- NodesTable: [NodeID: VARCHAR (PK), Latitude: FLOAT, Longitude: FLOAT, NodeCategory: VARCHAR]
- TransitionsTable: [RouteID: VARCHAR (PK), SourceNode: VARCHAR (FK), DestNode: VARCHAR (FK), BaseCost: FLOAT, TemporalRestrictionCode: INT]
- ConstraintsTable: [ConstraintID: VARCHAR, TargetNode: VARCHAR, ThrottlingFactor: FLOAT, StartTime: TIMESTAMP, EndTime: TIMESTAMP]

Environmental checks verify path compliance against these constraints in real-time, completely locally.

A.4 Prompt Engineering and Instruction Templates

Algorithm 1 utilizes the InjectConstraint module to dynamically insert failure details back into the planner model. Below is the precise instruction envelope used for executing this test-time alignment step:

```
=====
SYSTEM INSTRUCTION ENVELOPE: EXECUTION-AWARE PLAN REVISION
=====

You are an execution-aware trajectory correction model. An external predictive
critic has analyzed your previously generated long-horizon plan and detected
a critical environmental execution failure point.

[ORIGINAL HIGH-LEVEL GOAL]
{User_Goal_G}

[CURRENT PLAN CANDIDACY TRAJECTORY]
{Planned_Trajectory_Tau}
```

[DETECTED RISK LOCALIZATION HIGHLIGHT]

- Failed Step Index: {Step_t_Star}
- Action Element: {Action_Content}
- Environmental Boundary Contradiction: Highly probable structural collapse or unmodeled constraint conflict at this specific transition point.
- Critic Failure Confidence: {Failure_Probability}

[REVISION REQUIREMENTS]

1. Rewrite the trajectory sequence starting from the Failed Step Index.
2. Maintain all successful actions executed prior to the failure step index.
3. Select alternative routing parameters or package dependencies that bypass the flagged node without invalidating global goal bounds (temporal/financial).
4. Output your corrected plan using the identical structural JSON syntax.

Corrected Execution-Aware Plan Output:

=====