

Zen Agent Framework: Autonomous AI Systems

Technical Report v2025.07

Zen LM Research Team
research@zenlm.org

July 2025

Abstract

We present the Zen Agent Framework (ZAF), a production-grade system for building autonomous AI agents powered by Zen MoDE (Mixture of Distilled Experts) models. ZAF introduces three core innovations: hierarchical task decomposition that recursively breaks complex goals into executable subtasks, asynchronous tool calling that parallelizes independent actions, and context-adaptive compression that maintains long-horizon coherence within fixed context windows. On GAIA (all levels), ZAF achieves 72.4% accuracy. On SWE-bench Verified, ZAF resolves 48.3% of GitHub issues. On WebArena, task completion reaches 41.8%. We describe the framework architecture, tool protocol, multi-agent coordination, and memory systems in detail.

1 Introduction

Modern AI agents must autonomously plan, act, and adapt over extended horizons involving heterogeneous tools, dynamic environments, and incomplete information. While large language models provide the reasoning substrate, translating raw model capability into reliable agentic behavior requires careful engineering of the surrounding system.

The Zen Agent Framework addresses five fundamental challenges:

1. **Long-horizon planning:** Tasks requiring tens or hundreds of steps cannot be solved in a single model forward pass.
2. **Tool interoperability:** Agents must interface with diverse external systems—code executors, web browsers, databases, APIs—through a unified protocol.
3. **Context management:** Model context windows are finite; effective agents must compress, summarize, and selectively retrieve relevant history.
4. **Multi-agent coordination:** Complex tasks benefit from parallel execution by specialized subagents.
5. **Fault tolerance:** Real-world tool calls fail; agents must detect, diagnose, and recover from errors.

2 Agent Architecture

2.1 Core Loop

ZAF implements a ReAct-style observe-think-act loop with extensions for asynchronous execution and hierarchical task management:

Algorithm 1 Zen Agent Core Loop

```
1: goal  $\leftarrow$  user_input
2: plan  $\leftarrow$  HierarchicalPlanner(goal)
3: memory  $\leftarrow$  WorkingMemory.init()
4: while plan.incomplete() do
5:   task  $\leftarrow$  plan.next_executable()
6:   obs  $\leftarrow$  memory.relevant(task)
7:   thought, action  $\leftarrow$  ZenMoDE.reason(task, obs)
8:   if action.is_tool_call() then
9:     result  $\leftarrow$  ToolExecutor.run_async(action)
10:    memory.update(task, thought, action, result)
11:   else if action.is_subtask() then
12:     plan.expand(action.subtasks)
13:   else if action.is_answer() then
14:     plan.complete(task, action.answer)
15:   end if
16: end while
17: return plan.final_answer()
```

2.2 Hierarchical Task Decomposition

The planner decomposes goals recursively until tasks are atomic (single-step executable). The decomposition depth is bounded by $D_{\max} = 5$ to prevent infinite recursion. Each node in the task tree carries:

- **Goal:** Natural language description of the objective.
- **Dependencies:** Ordered set of prerequisite tasks.
- **Success criteria:** Verifiable conditions for task completion.
- **Resources:** Tools, permissions, and data required.

The decomposition scoring function ranks candidate subtask sets by estimated parallel executability:

$$\text{score}(\mathcal{S}) = \frac{|\mathcal{S}|_{\text{independent}}}{\text{depth}(\mathcal{S})} \cdot P(\text{success} \mid \mathcal{S}) \quad (1)$$

3 Tool Protocol

3.1 Tool Definition Standard

All ZAF tools conform to a standard definition schema inspired by the Model Context Protocol (MCP):

```
{
  "name": "code_execute",
  "description": "Execute Python code in isolated sandbox",
  "input_schema": {
    "type": "object",
    "properties": {
      "code": {"type": "string"},

```

```

    "timeout_seconds": {"type": "integer", "default": 30}
  },
  "required": ["code"]
},
"output_schema": {
  "type": "object",
  "properties": {
    "stdout": {"type": "string"},
    "stderr": {"type": "string"},
    "exit_code": {"type": "integer"}
  }
}
}
}

```

3.2 Async Tool Execution

Independent tool calls are dispatched concurrently. The dependency graph $G = (V, E)$ where V is the set of tool calls and E represents data dependencies enables optimal scheduling:

$$T_{\text{parallel}} = \max_{\text{chain} \in G} \sum_{v \in \text{chain}} t_v \quad (2)$$

compared to sequential execution $T_{\text{sequential}} = \sum_{v \in V} t_v$.

For typical agent workloads with parallelism factor 3.4, async execution reduces wall-clock time by 68%.

3.3 Tool Categories

Table 1: ZAF built-in tool categories and instances

Category	Tools	Avg Latency
Code execution	Python, JavaScript, Bash	1.2s
Web	Browser, fetch, search	2.8s
File system	Read, write, list, diff	0.08s
Version control	git clone/diff/commit	0.9s
Database	SQL query, schema inspect	0.3s
API	REST, GraphQL, MCP clients	1.4s
Computation	Math, data analysis	0.6s
Communication	Email, Slack, GitHub Issues	1.1s

4 Memory Systems

4.1 Working Memory

Working memory holds the current task trajectory—observations, thoughts, and actions—within the model’s context window. As context grows, ZAF applies hierarchical compression:

$$C_{\text{compressed}} = \text{Summarize}(C_{\text{old}}, \tau_{\text{compression}} = 0.4) \quad (3)$$

Compression is triggered when context utilization exceeds 75%, preserving the most recent 25% verbatim and summarizing the rest.

4.2 Episodic Memory

Completed task trajectories are stored in a vector database as episodic memories. On new tasks, the agent retrieves relevant episodes:

$$\text{episodes} = \text{TopK}(\{e_i : \text{sim}(\mathbf{e}_{\text{task}}, \mathbf{e}_{e_i}) > \theta\}, K = 5) \quad (4)$$

Retrieved episodes are summarized and prepended to context as “prior experience”, reducing exploration cost on similar tasks by 34%.

4.3 Semantic Memory

Factual knowledge from completed tasks is structured into a knowledge graph updated via:

$$\mathcal{G}_{t+1} = \mathcal{G}_t \cup \text{Extract}(\text{trajectory}_t) \quad (5)$$

Knowledge graph queries provide instant access to facts discovered in prior sessions without re-exploration.

5 Multi-Agent Coordination

5.1 Agent Roles

ZAF supports three agent roles in multi-agent configurations:

- **Orchestrator:** Decomposes the high-level goal, assigns subtasks to workers, and synthesizes results.
- **Worker:** Executes assigned subtasks with access to a specific tool subset.
- **Critic:** Reviews worker outputs for correctness, calling for revision when quality thresholds are not met.

5.2 Communication Protocol

Agents communicate via structured messages over a shared event bus. Each message carries:

- Sender and recipient agent IDs.
- Task ID for correlation.
- Message type: `ASSIGN`, `RESULT`, `CRITIQUE`, `ESCALATE`.
- Payload: task description or result artifact.
- Priority: 0–10 for scheduling.

5.3 Consensus and Verification

For high-stakes tasks (e.g., code deployment, financial actions), ZAF employs multi-agent voting. Three independent worker agents execute the task, and the orchestrator accepts the result only if at least two agents agree:

$$\text{accept}(\text{result}) = \mathbf{1} \left[\sum_{i=1}^3 \mathbf{1}[\text{result}_i = \text{result}] \geq 2 \right] \quad (6)$$

6 Context Compression

6.1 Adaptive Summarization

ZAF uses Zen MoDE itself for context compression, applying a hierarchical summarization strategy:

1. **Tool result compression:** Raw tool outputs (e.g., HTML pages, log files) are summarized to their relevant portions before entering context.
2. **Step-level summaries:** Every 10 completed steps, prior steps are summarized into a condensed trajectory record.
3. **Episode summaries:** When a task tree branch completes, the full subtask trajectory is compressed to a paragraph.

Compression ratio is $6.8\times$ on typical agentic trajectories while preserving 94.2% of decision-relevant information (measured by downstream task accuracy with vs. without compression).

7 Benchmark Results

7.1 GAIA

GAIA evaluates real-world question answering requiring multi-step tool use. Questions range from Level 1 (single-tool) to Level 3 (complex multi-step with >10 tool calls).

Table 2: GAIA benchmark results

System	L1	L2	L3	Overall	Avg Steps
GPT-4o (ReAct)	74.8	41.2	18.4	52.4	6.2
Claude 3.7 (Extended)	78.4	46.8	22.1	56.8	7.8
ZAF (72B)	81.2	51.4	26.8	60.4	8.4
ZAF (236B)	86.4	58.2	34.1	67.2	9.1
ZAF (480B)	89.8	63.4	39.2	72.4	9.8

7.2 SWE-bench Verified

SWE-bench Verified evaluates agents on real GitHub issues from Python repositories.

Table 3: SWE-bench Verified resolution rate (%)

System	Resolved	Patch Applied	Tests Pass	Avg Edits
ZAF (72B)	38.4	52.1	71.2	3.8
ZAF (236B)	44.1	58.4	77.8	4.2
ZAF (480B)	48.3	63.2	82.4	4.6
ZAF (480B) + multi-agent	51.8	66.8	84.1	5.2

7.3 WebArena

WebArena tests web navigation agents across shopping, GitLab, Reddit, CMS, and mapping tasks.

Table 4: WebArena task success rate (%)

System	Shopping	GitLab	Reddit	CMS	Map	Overall
ZAF (72B)	32.4	28.1	34.8	31.2	44.8	33.2
ZAF (236B)	38.1	34.2	40.1	37.4	51.2	38.8
ZAF (480B)	42.8	38.4	44.2	41.8	56.4	41.8

7.4 AgentBench

Table 5: AgentBench overall score across 8 environments

System	Score	Avg Turns
ZAF (72B)	4.82	8.4
ZAF (236B)	5.64	9.2
ZAF (480B)	6.18	9.8

8 Safety and Guardrails

8.1 Action Verification

All tool calls are classified by risk level before execution:

- **Safe:** Read-only operations, executed immediately.
- **Moderate:** Write operations to sandboxed environments, executed with logging.
- **High:** External API calls with financial or irreversible effects, require confirmation.
- **Blocked:** Actions matching blocklist patterns (e.g., credential exfiltration) are refused.

8.2 Trajectory Auditing

All agent trajectories are logged with cryptographic hashes linking each action to its preceding context. This enables post-hoc audit of agent behavior and detection of manipulation attempts.

9 Conclusion

The Zen Agent Framework demonstrates that principled engineering of the agentic loop—hierarchical planning, async tool execution, adaptive memory, and multi-agent coordination—substantially improves performance over naive ReAct baselines. ZAF achieves 72.4% on GAIA and 48.3% SWE-bench Verified resolution, establishing Zen MoDE as a leading backbone for autonomous AI systems.

References

- [1] Mialon, G. et al. GAIA: A Benchmark for General AI Assistants. *ICLR*, 2024.
- [2] Jimenez, C. et al. SWE-bench: Can Language Models Resolve Real-world Github Issues? *ICLR*, 2024.
- [3] Zhou, S. et al. WebArena: A Realistic Web Environment for Building Autonomous Agents. *ICLR*, 2024.

- [4] Yao, S. et al. ReAct: Synergizing Reasoning and Acting in Language Models. *ICLR*, 2023.
- [5] Liu, X. et al. AgentBench: Evaluating LLMs as Agents. *ICLR*, 2024.