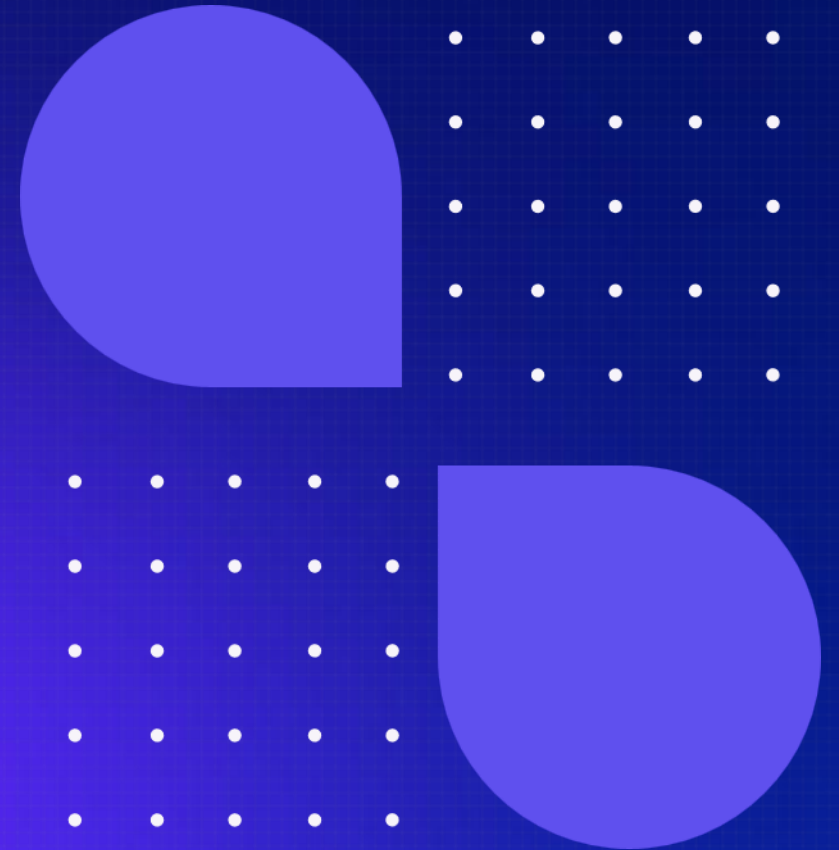
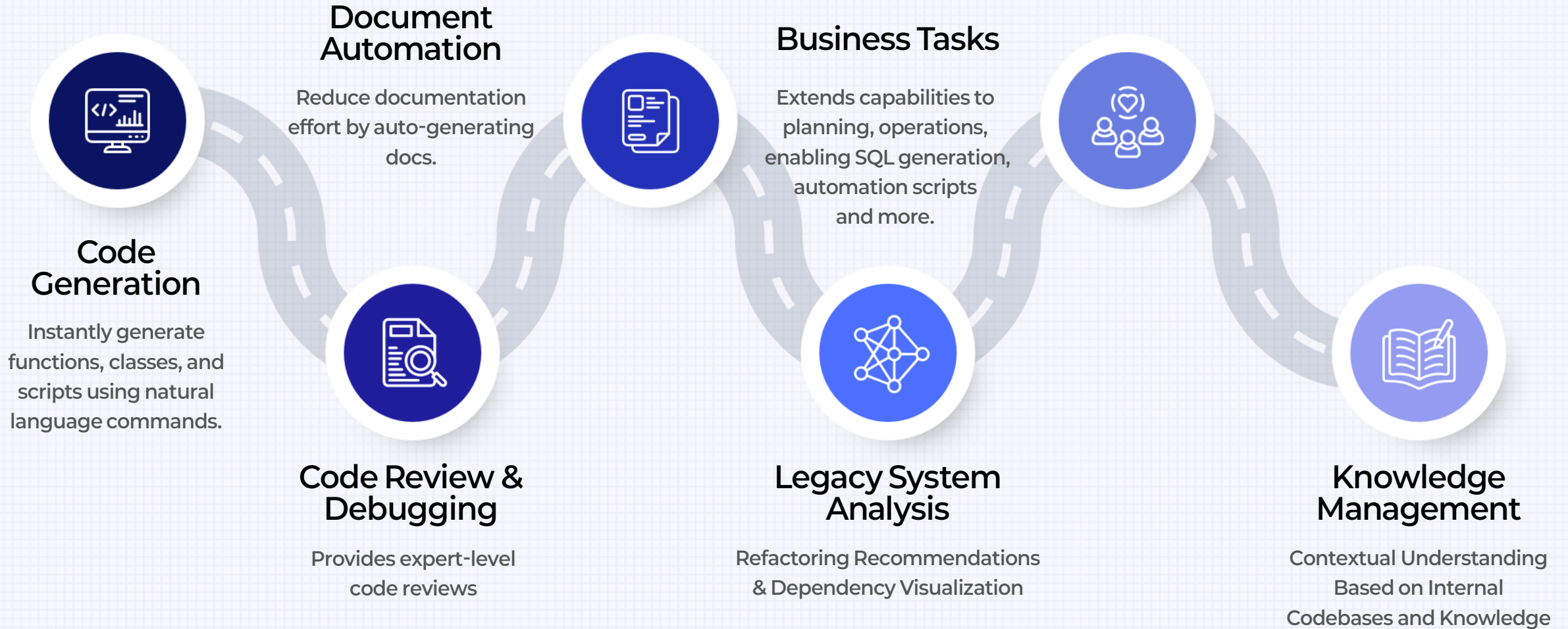


Athena Code Assistant

Provides Security and efficiency in
On-Premise environment



| AI Tools for the Entire Organization, Beyond Coding



| Why Existing AI Code Assistants Fall Short in Enterprise Environments

Blocking Access to External AI Services Within Secure Networks

Cloud-based AI code assistants require external API access, making them unsuitable for secure and network-segregated environments.

Inability to Control AI Usage at the Team and Organizational Levels

Personal SaaS tools lack organization-wide access control, usage limits, and AI usage visibility.

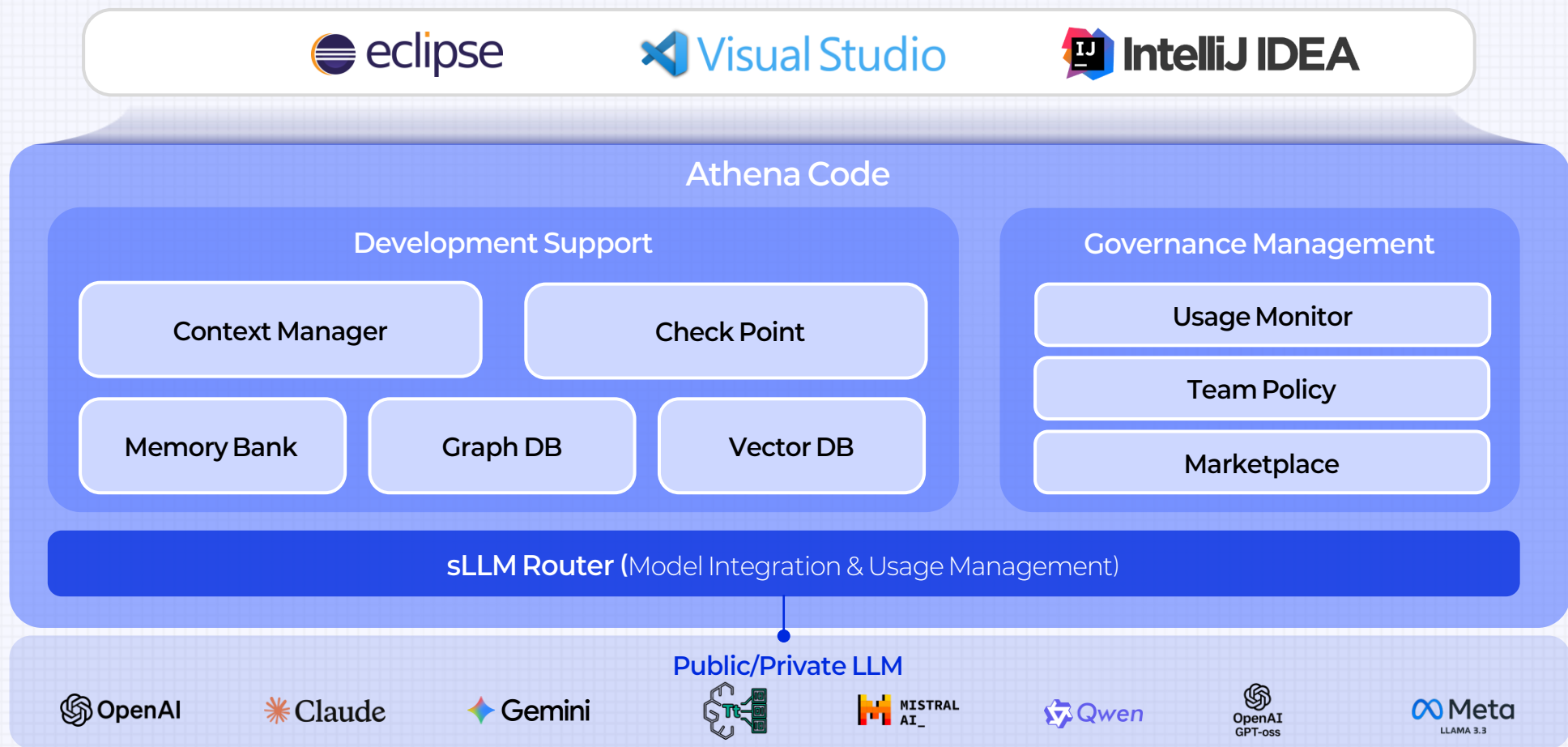
Difficulty Ensuring the Quality and Security Reliability of AI-Generated Code

AI-generated code requires quality and security validation, with auditability essential in financial and public-sector environments.

Limitations in Understanding Large-Scale Legacy Codebases

File-level AI tools lack system-wide context, limiting their understanding of complex dependencies and code impact.

Athena Code Assistant for Enterprise Organizations



Overcoming the Limitations of Existing AI Code Assistants



Enhanced Security

Isolated Network Operation

FOCUS

- Independent Internal Network Operation
- Air-Gapped Ready

Technologies Applied

- Selective Public/Private LLM
- sLM Router-Based Model Management

Automated Audit Compliance

FOCUS

- Automated AI Usage Logging
- Admin Dashboard Management & Monitoring

Technologies Applied

- Automated Audit Logging
- Team Policy-Based Access Control



Strengthened Development Capabilities

Contextual Understanding of Legacy Code

FOCUS

- Unified System Architecture
- Instant Impact Analysis

Technologies Applied

- Graph DB Visualization
- Memory Bank & Vector DB Knowledge Search

Visibility into AI Usage

FOCUS

- Real-Time Usage Monitoring
- User-Based Model Access

Technologies Applied

- Usage Monitor Dashboard
- User-Based Model Usage Limits

Develop Faster, Stay Safer.

01

Automation of Repetitive Tasks

Context-Persistent Development Environment

Context

CheckPoint

/ Command

02

Legacy & Large-Scale Code Analysis

Instant Structural Analysis of Large-Scale Codebases

Graph DB

Memory Bank

03

Security & Audit Compliance

Complete AI Usage Tracking & Auditability

Log

Compliance

04

Cost & Usage Control

Real-Time Cost Visibility by Team & User

sLLM Router

AI Monitoring

Usage

05

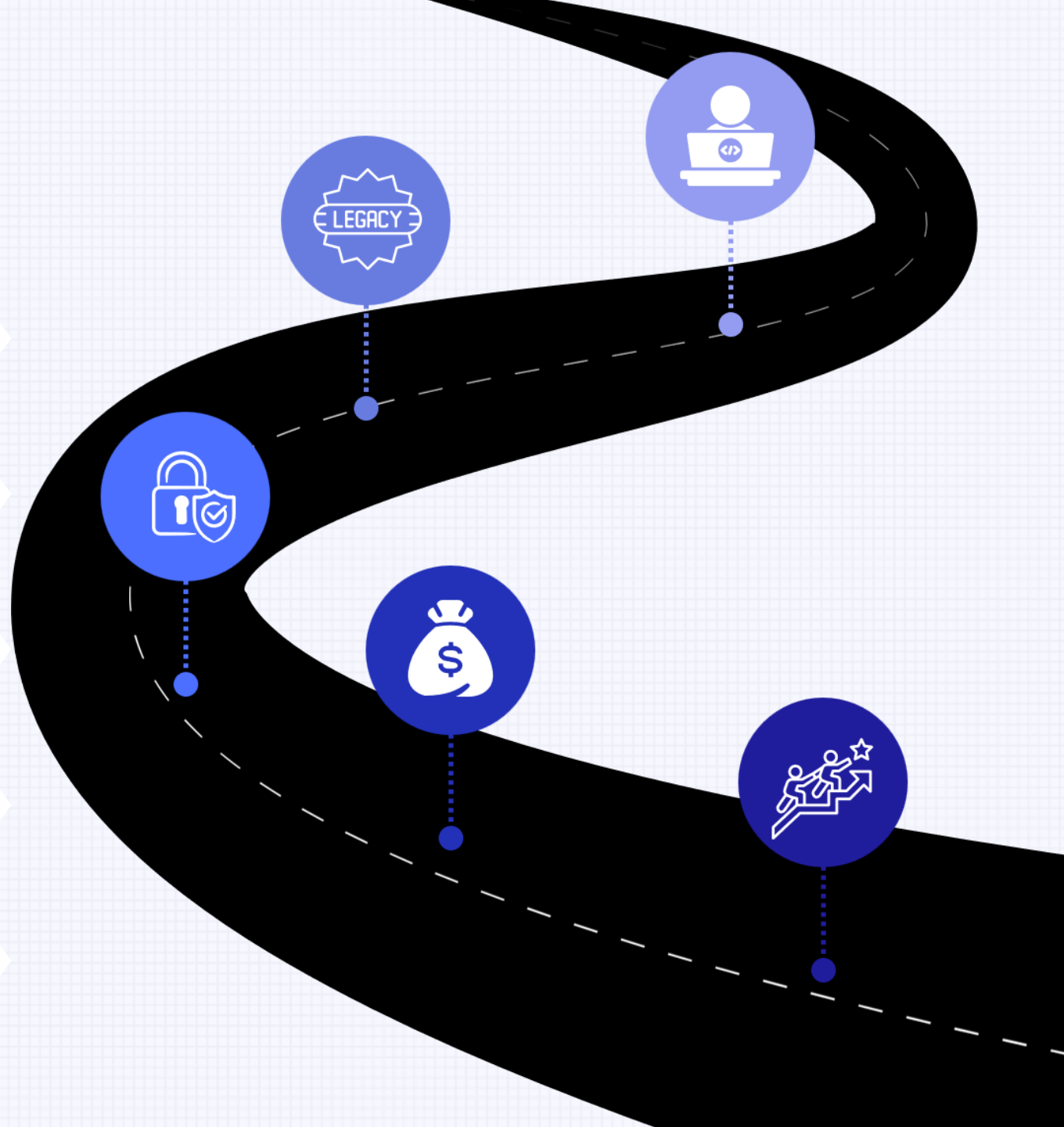
Organization-Specific Customization

Organization-Specific AI Environment Configuration

Rules

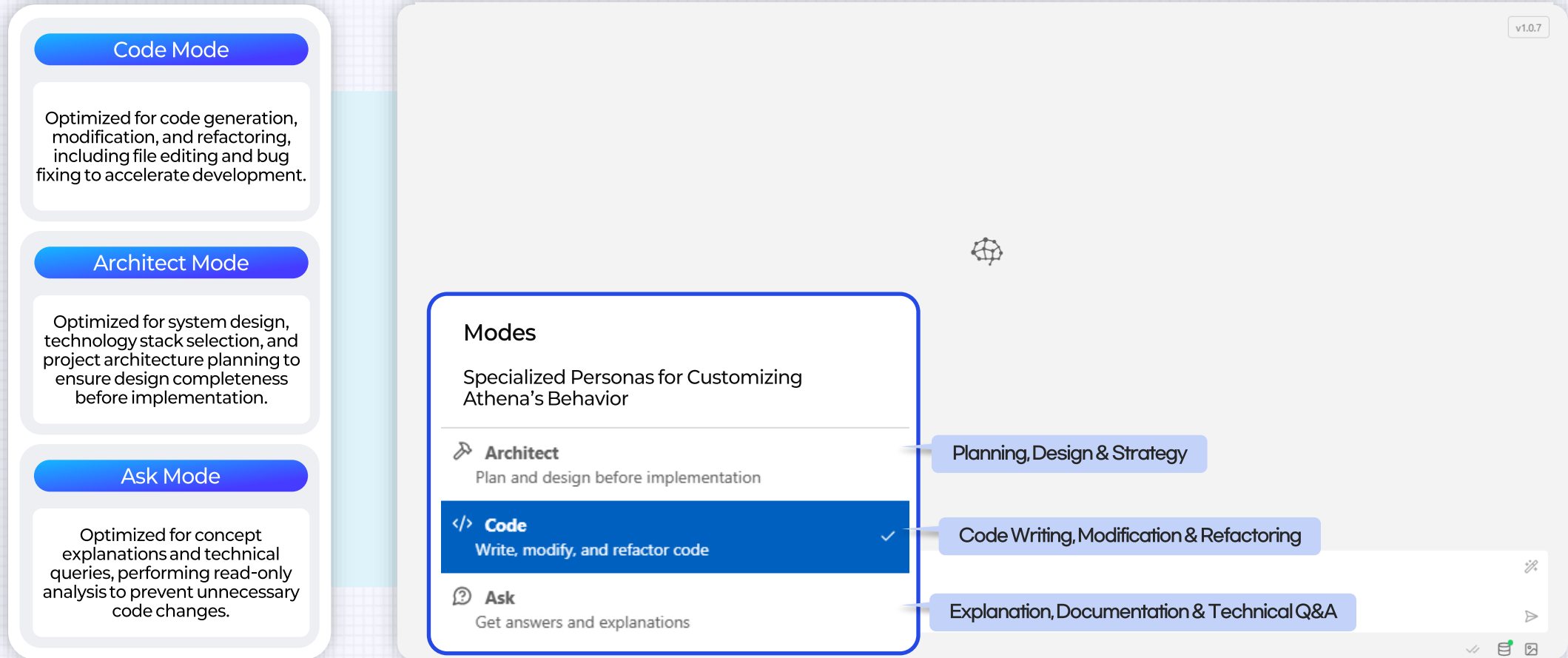
MCP

Market Place



Modes

Define AI's scope of operation by task to prevent unnecessary file access and code changes.



Context Setting

Controls the information included in the AI's context window, influencing token usage and response quality.

Context Adjustment

Adjusting open tabs, workspace files, and concurrent read limits to optimize cost, response quality, and overall performance

Automatic Blocking

Automatically excluding sensitive files and directories from AI access and search to ensure transparent security governance

Ensuring Task Stability

Automatically summarizing and compressing previous conversations when the context usage threshold is reached, enabling uninterrupted task continuity

Direct Context Scope Adjustment

Providing only the necessary information to the AI

Exclusion of Sensitive Files from AI Reference

Blocking security-sensitive files and confidential code

Automatic Context Compression

Maintaining response quality and reducing token costs in long conversations

Open Tab Context Limit

20

Workspace File Context Limit

200

Concurrent File Read Limit

5

☒ Show Files in Lists and Search

Automatic File Read Reduction Threshold

-1 Lines ☒ Always Read Full Files

☒ Intelligent Context Compression Auto-Trigger

✱ Compression Trigger Threshold

Global Defaults

90%

Check Points

Automatically creates a snapshot before modifying files, allowing you to safely restore the previous state at any time.

Auto-Snap Shot

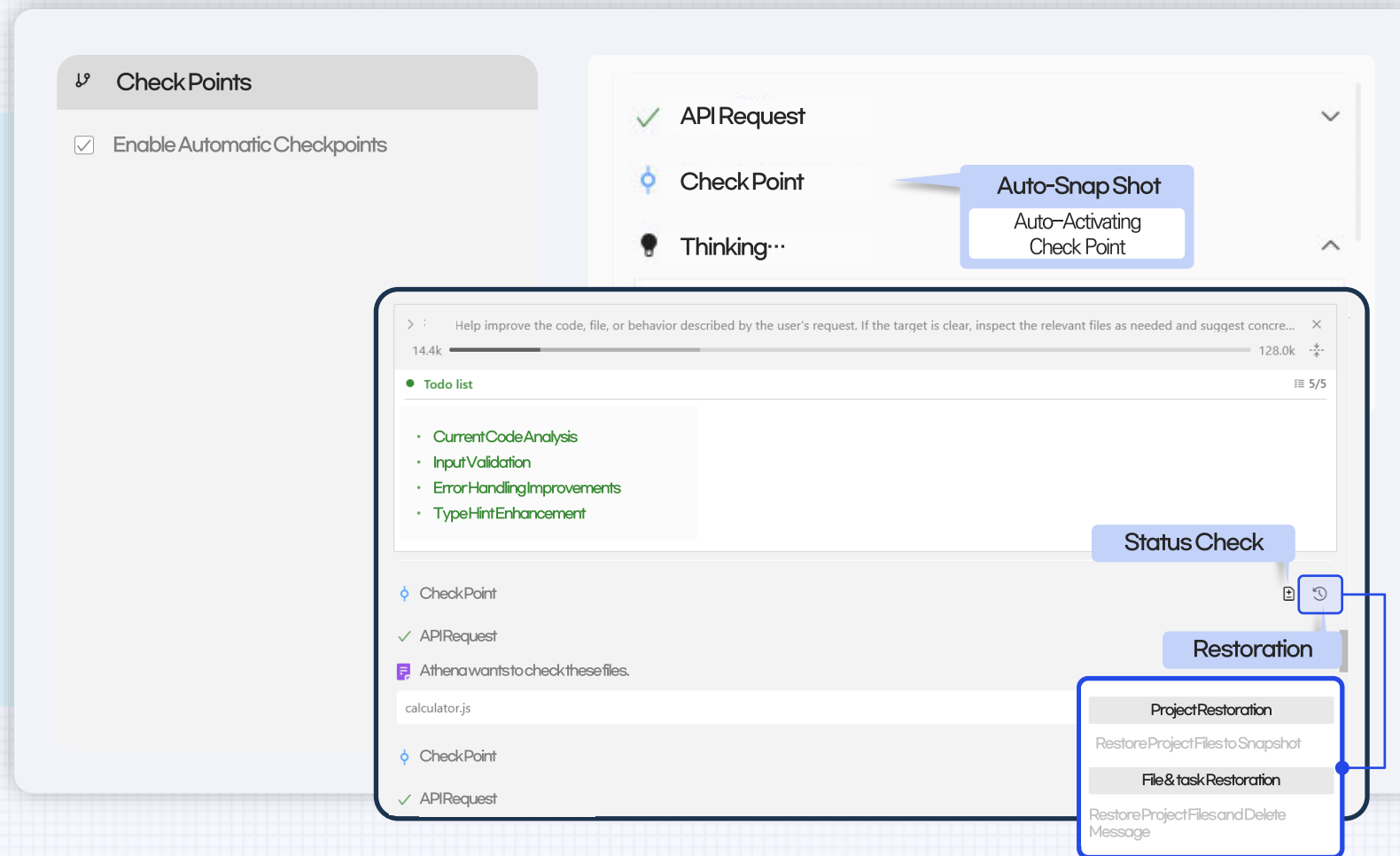
Automatically creates Git-based snapshots before file modifications begin, preserving the complete change history.

State Restoration

Enables restoration to a checkpoint at any desired point, preventing data loss caused by mistakes or unintended changes.

Workflow Stability

Automatically managed during work without additional configuration, maintaining a stable development environment without interrupting the workflow.



Context Mentions & Slash Commands

Use @mentions to specify the resources the AI should reference, and execute repetitive tasks instantly with commands.

Context Source

Reference files, folders, terminals, Git, and URLs with @mentions to provide information to the AI.

Quick Execution

Register frequently used task patterns as slash (/) commands for instant execution without writing prompts.

Shared Standard Code

Share team-defined Rules and Commands through a marketplace and apply them instantly across all team members' IDEs.

Slash Commands

Manage slash commands to quickly execute custom workflows and tasks.

Slash Commands

Manageslashcommandstoquicklyexecutecustomworkflowsandtasks.

Commands

/init

/create-command
CommandCreation

/create-rules
RuleCreation

/improve
ImprovementSuggestions

/explain
Explanations&Comments

```
push_test_data.json
push_test_data.csv ·
push_test_data_multiplex.csv - D
```



Instant Context
Reference Linking

Morpheus MADP Client | URACLE

@

Instant Execution of
Custom
Commands

Memory Bank

Structures Important Decisions, Tasks, and Patterns to maintain consistent understanding across sessions.

Project Context Preservation

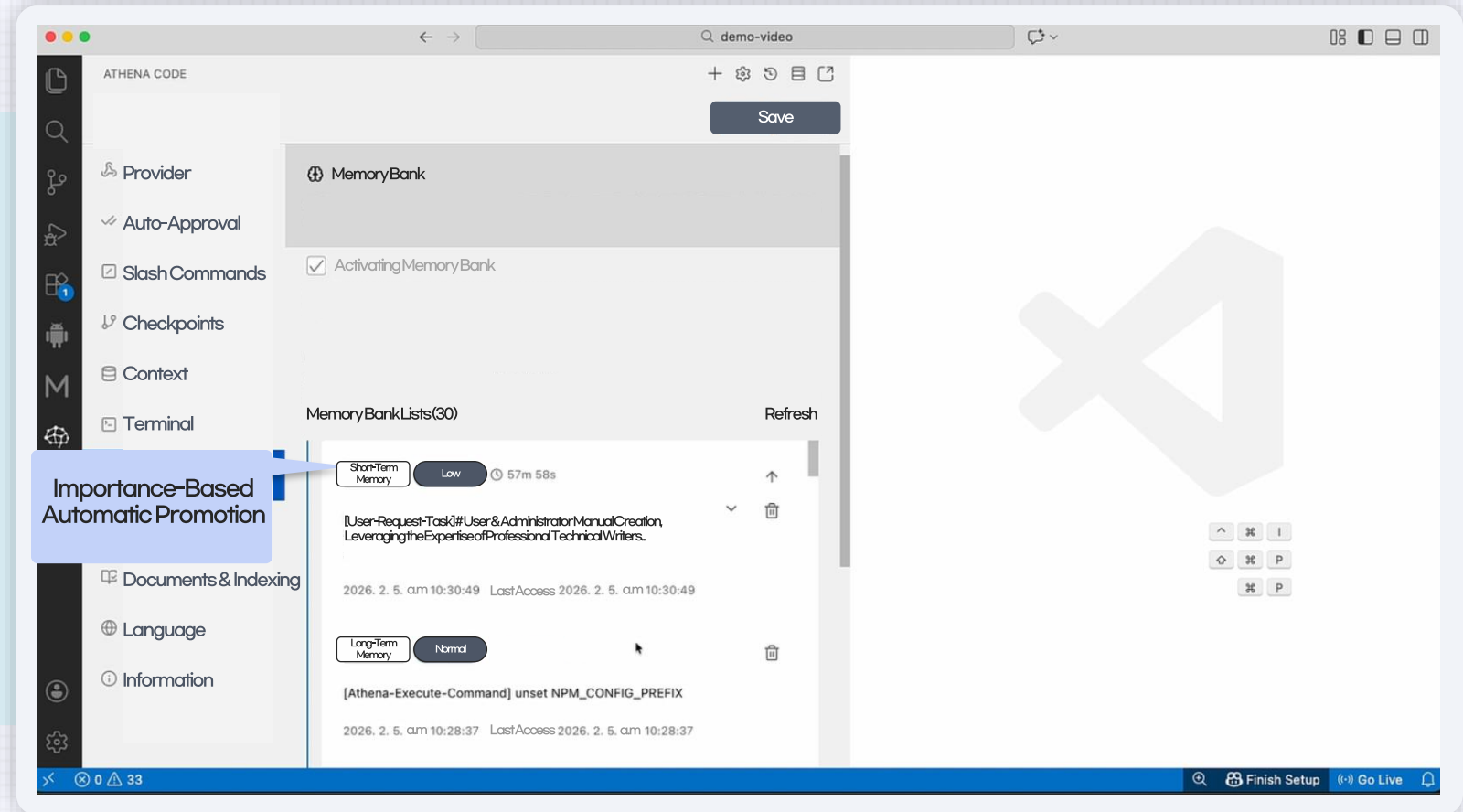
Preserve design decisions and work history across sessions using a hierarchical STM (Short-Term Memory) and LTM (Long-Term Memory) structure.

Non-Redundant Knowledge

Automatically detect duplicates based on vector similarity and reconcile conflicting information with existing content to maintain up-to-date knowledge integrity.

Instant Context Retrieval

Instantly retrieve relevant context from project-specific, independent knowledge bases to maintain continuity without interrupting the workflow.



Graph DB-Based Code Analysis

Visualize codebase call relationships and dependencies to quickly understand even large-scale legacy code.

Code Relationship Analysis

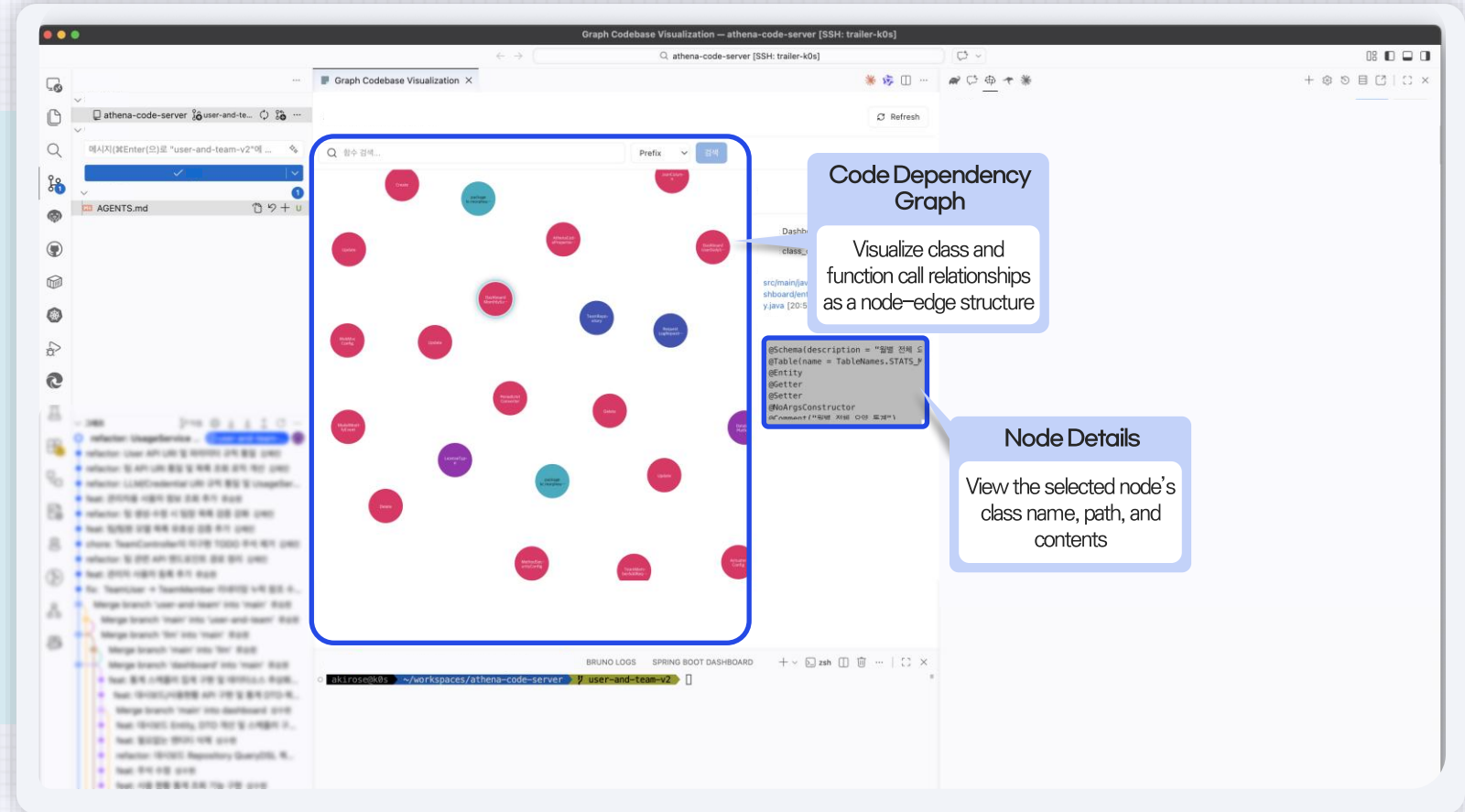
Visualize class and function call relationships using a graph database to understand the overall code structure.

Semantic Code Chunking

Automatically split code into function- and class-level units to enable precise code retrieval through natural language queries.

Legacy Code Understanding

Automatically index large-scale legacy code using dependency graphs, significantly reducing the time required for newly assigned developers to understand the codebase.



Track AI usage history across the entire team in real time, with the ability to view and export detailed information for each request at any time.

History Review

Record and retrieve all AI request history in real time based on date/time, team, user, model, tokens, cost, and status.

Request History Review

Review the full Request, Response, and Trackback details of individual requests to assess response quality and identify anomalies.

Audit Compliance

Filter logs by period, team, user, and model, then export them for internal audits and compliance requirements.

HistoryLog

View All Request History

Filter and display request history by date/time, user, tokens, and cost.

2026.05.21 ~ 90 Days

Name	Team	Model	Status	Token	Cost(\$)	View Details	
Results 219							
Log	Team	Name		Token	Cost(\$)	Status	View Details
2026-05-28 08:49	-	yongbum7186@uracle.co.kr	qwen3	53,641	0.016	Success	Details
2026-05-28 08:48	-	yongbum7186@uracle.co.kr	qwen3	51,888	0.014	Success	Details
2026-05-28 08:48	-	yongbum7186@uracle.co.kr	qwen3	11,889	0.004	Success	Details
2026-05-28 08:48	-	yongbum7186@uracle.co.kr	qwen3	10,021	0.003	Success	Details
2026-05-28 08:48	-	yongbum7186@uracle.co.kr	qwen3	27,792	0.008	Cancel	Details
2026-05-28 08:48	-	yongbum7186@uracle.co.kr	qwen3	9,886	0.003	Success	Details
2026-05-27 15:46	-	yongbum7186@uracle.co.kr	qwen3	57,570	0.017	Success	Details
2026-05-27 15:45	-	yongbum7186@uracle.co.kr	qwen3	51,035	0.015	Success	Details
2026-05-27 15:45	-	yongbum7186@uracle.co.kr	qwen3	31,103	0.009	Success	Details
2026-05-27 15:45	-	yongbum7186@uracle.co.kr	qwen3	30,484	0.009	Success	Details

ID: chatlog-4470c741-c70f-470e-a27f-ff6a8861a947

Model: qwen3.5-35b

Model ID: 49

Start Time: 2026-05-28 08:49:52

End Time: 2026-05-28 08:50:01

Tokens: 53,641

Cost: 0

Request

Response

Trackback

AI Usage Visibility & Cost Control

Monitor AI usage and costs in real time by team, user, budget, and model to control AI usage across the organization.

Usage Monitoring

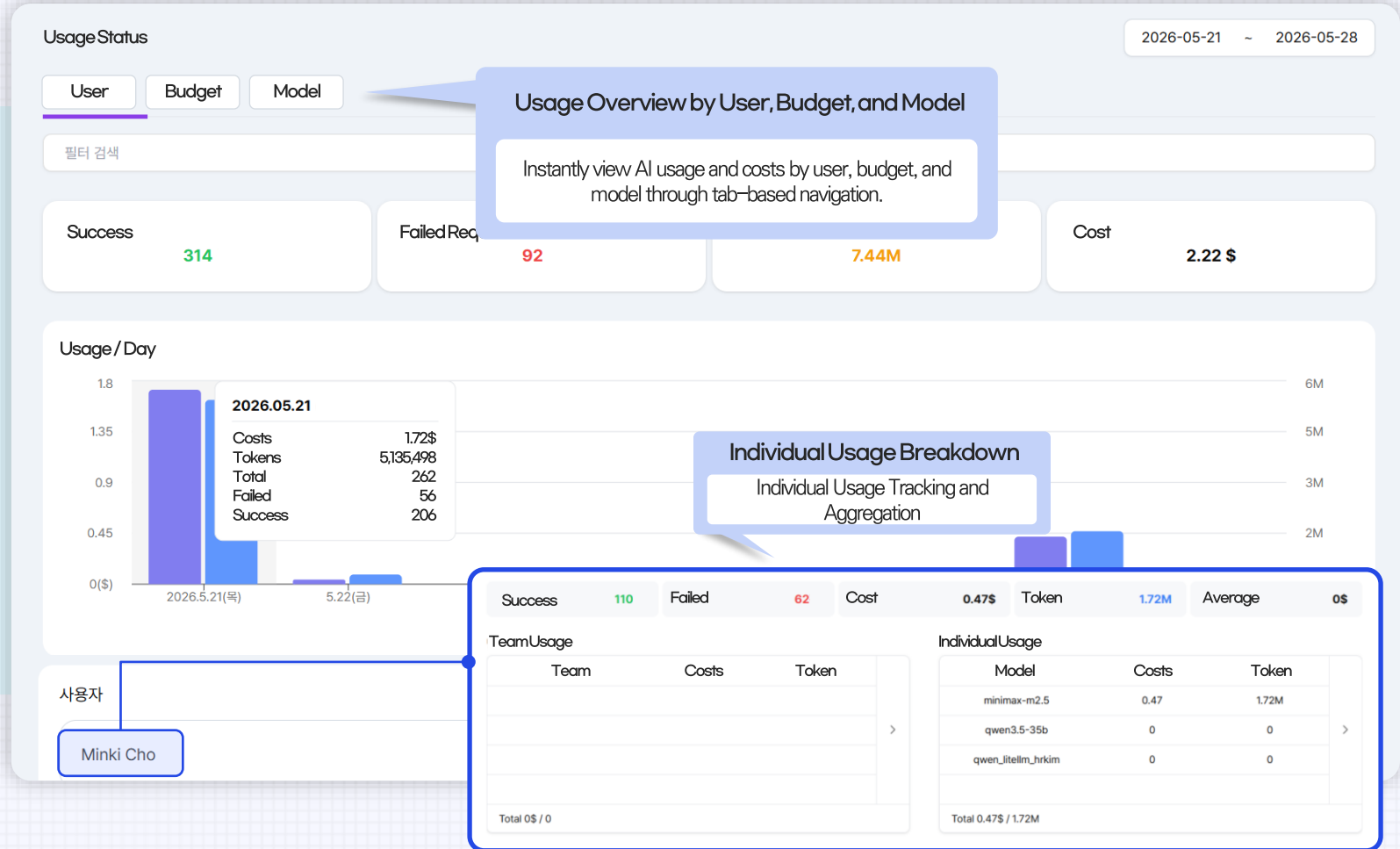
Real-Time Aggregation of Successful/Failed Requests, Tokens, and Costs for Immediate Visibility into Organization-Wide AI Usage.

User Analytics

Identify cost-intensive areas and detect anomalous usage through dedicated tabs for users, budgets, and models.

Cost Configuration

Set spending limits by user and team to proactively block usage before exceeding budgets and prevent unnecessary costs.



Rules Define the team's coding standards in a document, and AI automatically applies them across every developer's IDE.

Standard Distribution

Apply rules defined by team leads and senior developers instantly across all team members' IDEs to ensure consistent code quality.

Rule Enforcement

Detect rule violations in real time during coding and provide automatic fix suggestions to minimize code review overhead.

Project Segmentation

Manage standards flexibly by separating organization-wide rules from project-specific rules.

ATHENA CODE

> Task: Create a user login function based on these rules.

15.4k

128.0k

Provided Features

Function	Description
<code>registerUser()</code>	User Registration
<code>authenticateUser()</code>	User Login Authentication
<code>removeUser()</code>	User Deletion
<code>getAllUsers()</code>	Retrieve All Users

Team Coding Standards Compliance

- ✓ Prohibit console.log; always use logger.info
- ✓ Function names must always begin with a verb
- ✓ Variable names must use camelCase (Came1 Case → camelCase)

Team Coding Rules Setting

- Prohibit console.log; always use logger.info
- Function names must always begin with a verb
- Variable names must use camelCase (Came1 Case → camelCase)

```
const logger = {  
  info: (message) => console.log(`[INFO] ${message}`),  
  error: (message) => console.error(`[ERROR] ${message}`),  
};
```

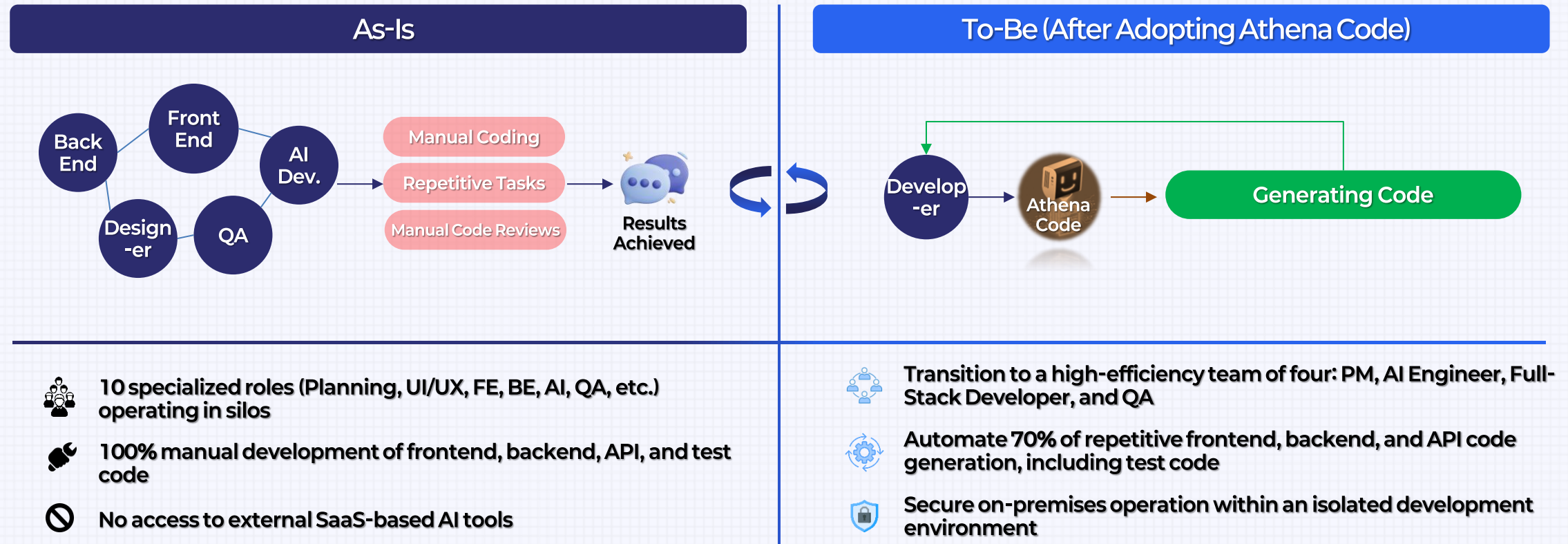
Key Feature Comparison

Each product is optimized for different environments and purposes. This reference summarizes the support status of key features based on publicly available information.

※ As of June 2026

Features	 CURSOR	 Claude	 Copilot	Athena Code
Code Autocomplete	✓	✓	✓	✓
Full Codebase Context Understanding	✓	✓	✓	✓
IDE Support (VS Code)	✓	✓	✓	✓
IDE Support (IntelliJ · JetBrains)	✓	✓	✓	✓
IDE Support (Eclipse)	✗	✗	✓	✓
Long-Term Context Preservation (Across Sessions)	✓	✓	✗	✓
Fully Closed-Network & Air-Gapped Operation	✗	✗	✗	✓
On-Premises Deployment	△	✗	✗	✓
Real-Time Team Usage & Cost Monitoring	△	△	△	✓
Automated Audit Logging & Export	△	△	△	✓
Team Coding Standards Synchronization	△	✗	✗	✓
Graph DB-Based Code Structure Visualization	△ Requires separate installation and integration of external tools			✓

Internal Chatbot Development Use Case



Based on internal PoC results, approximately 40% reduction in development effort and 50% reduction in development time.

Legacy Web Modernization Use Case



Significantly reduce the cost, timeline, and risks of legacy modernization within an isolated development environment.

Support for Code Assistant Environments of Various Scales

Supports not only large-scale environments but also small- and medium-sized environments with limited GPU resources.

※ Qwen3.6-35B-A3B sLM

Category	Device	Available Memory	AI Model Memory Usage	Recommended Number of Users	Key Features
Personal	AMD AI Max 395+	128GB (Shared)	~20GB	1	High-Performance Laptop
	MacBook M5 Pro	64GB (Shared)			Personal Laptop Environment Optimization
Team	Mac Studio M3 Ultra	96GB (Shared)		3	Small Team Shared Use, Multi-Session Optimization
	RTX PRO 6000	96GB VRAM		6~10	Medium Team Operations Optimization, Maximum Performance

💡 The Code Assistant does not consume GPU resources while users are typing. Actual GPU utilization is only around 10~20% of the total, meaning a single RTX PRO 6000 can effectively support a team of 6~10 users.

Thankyou

H. uracle.co.kr

E. sales@uracle.co.kr

T. +82-2-3479-4400