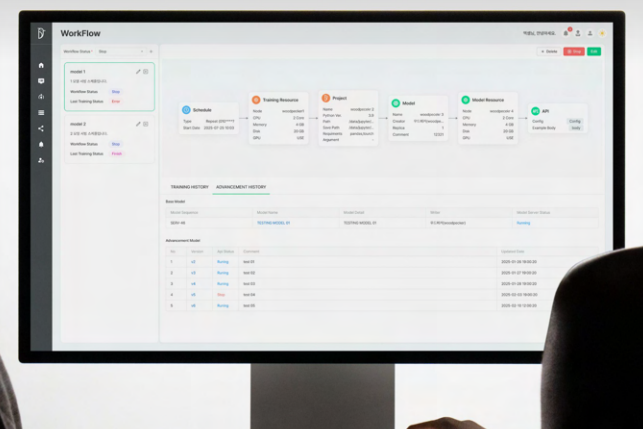




Big Data Analytics. AI Model Development. Anomaly Detection.

Woodpecker - AI Model Development



XAI Ops - Anomaly Detection



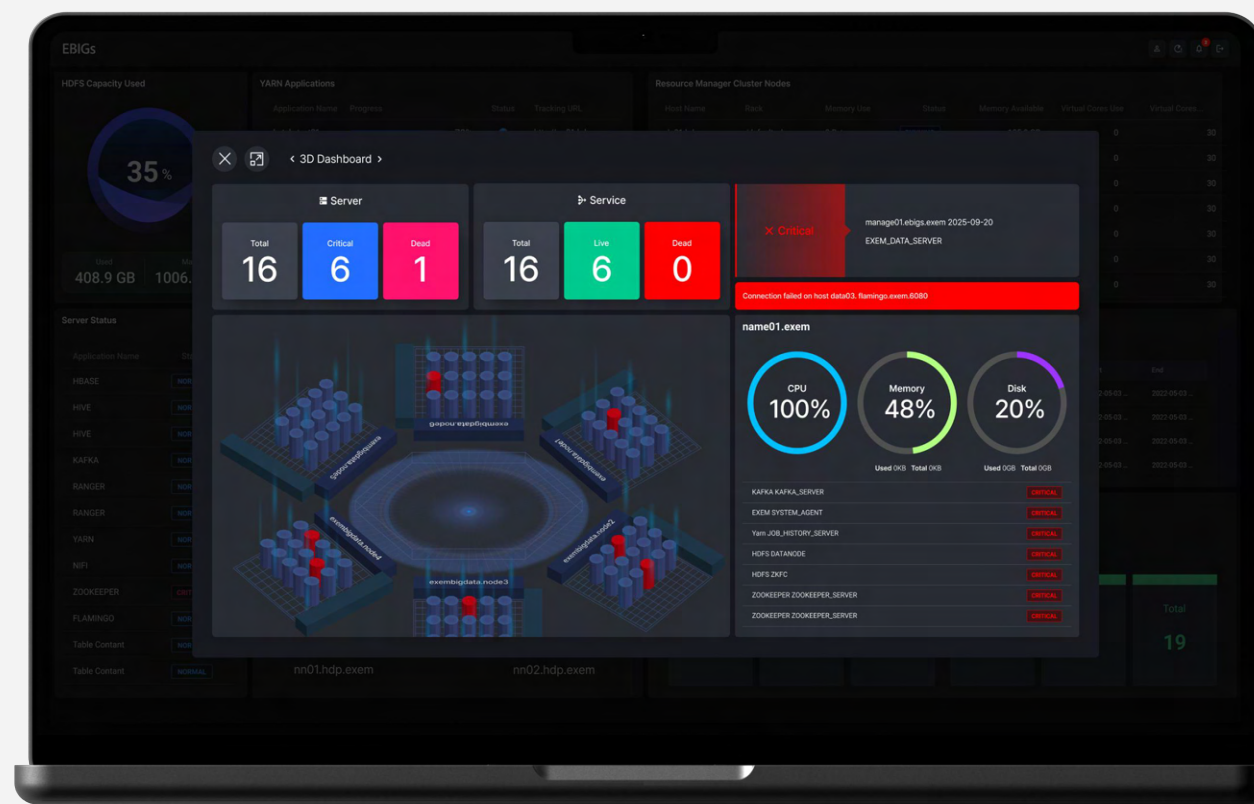
EBIGs - Big Data Analytics

Comprehensive AI Solutions Spanning the Entire Spectrum

EBIGs is a big data operations platform based on the Apache Hadoop Ecosystem that helps reliably store, process, and analyze big data.

Woodpecker is a Kubernetes based self-service AI analytics platform that enables anyone to easily process big data and perform AI analysis.

XAI Ops is an AIOps solution that analyzes server, DB, network, and log data with AI to instantly detect anomalies and predict incidents, maximizing operational efficiency and business stability.

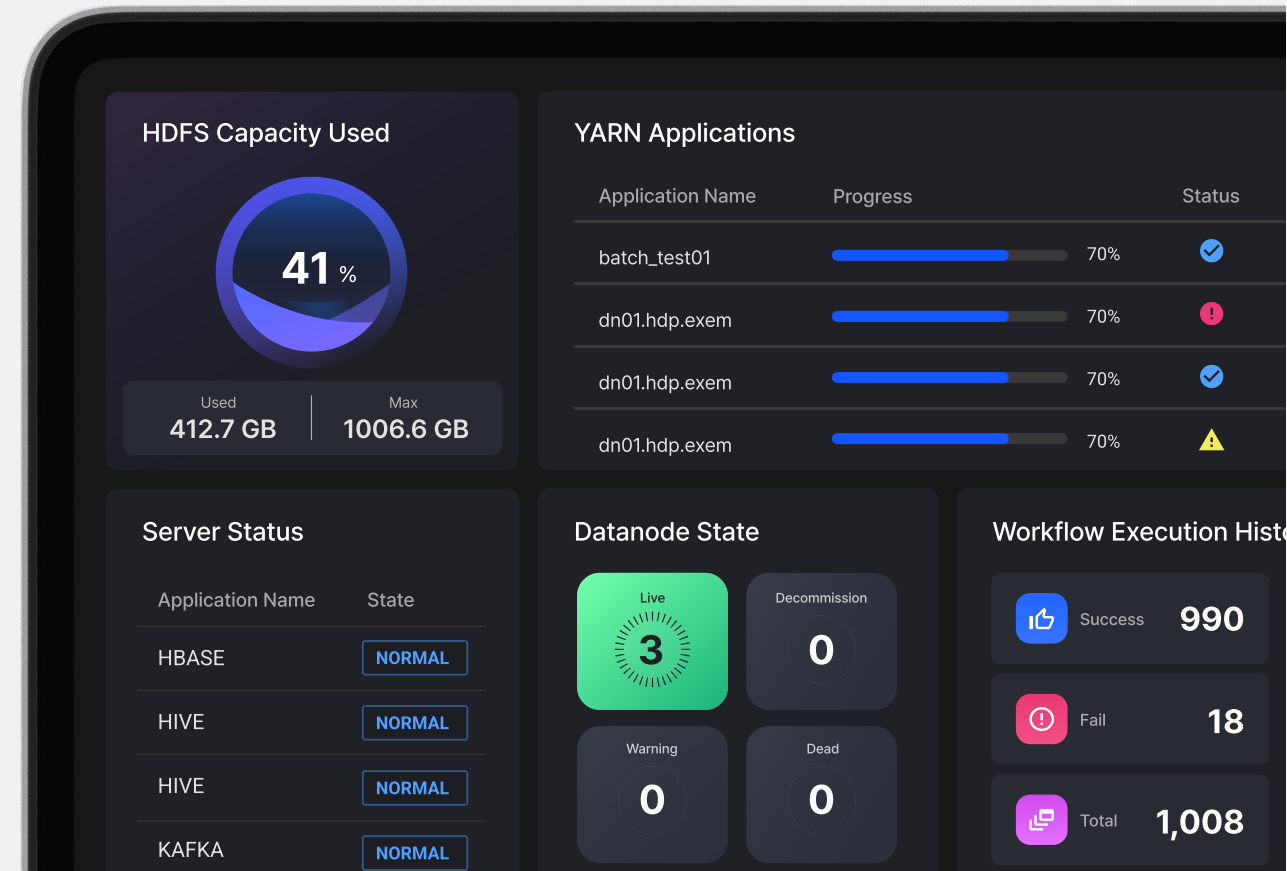




Why EBIGs Stands Out

Unified Management Platform from Big Data Configuration to Operations

EBIGs is a big data operations platform based on the Apache Hadoop Ecosystem. Provide compatibility verified open-source stacks and dedicated operations consoles to standardize configuration, security, and operations issues for safe, consistent monitoring and operations. Reduce operational burden while achieving both stability improvement and cost savings.



Customer Stories

Public Institution Regional Integrated Data Hub	Built an Apache Hadoop based big data platform for systematic regional data integration and utilization. Adopted EBIGs to unify regionally distributed data, improving work efficiency and establishing a foundation for data driven policy making.
Public Institution Disaster Recovery System	Implemented a disaster recovery system(DRS) using EBIGs for non stop operation of core systems. Built an environment for rapid service recovery during disasters through real time data replication, thoroughly ensuring data integrity and consistency. Secured both business continuity and customer trust.

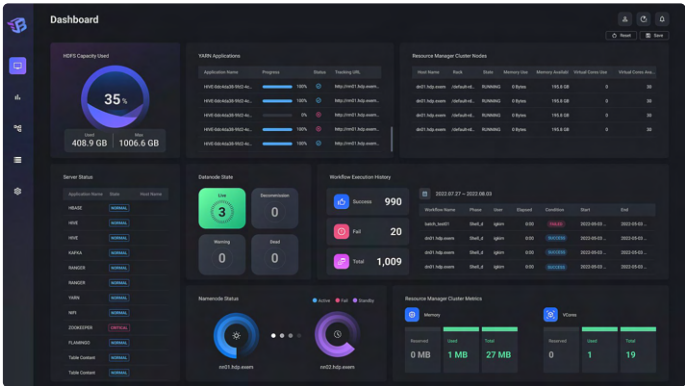
Product Highlights

 All Data Summarized in One Screen	 Real Time Data Insights Driving Business Decisions
Visualize the Hadoop ecosystem, a leading big data management framework, in an easy to understand format. Check server, service, and node status at a glance for efficient data management.	Collect and process massive data in real time to deliver key insights quickly and accurately. Save database analysis time and focus on critical decision making.

1 Hadoop Cluster Unified Monitoring Dashboard

Entire cluster status in a single screen

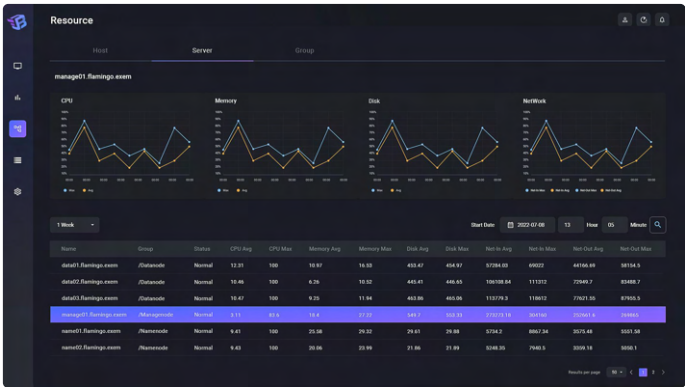
- Verify server, service, node status and key metrics in real time (HDFS, YARN, Hive, etc.)
- Analyze status, capacity, and resource trends for components: Resource Manager, Name Node, Node Manager
- Instantly identify topology and anomaly points via 3D node map



2 Service/System Resource

Real time verification of cluster node resource

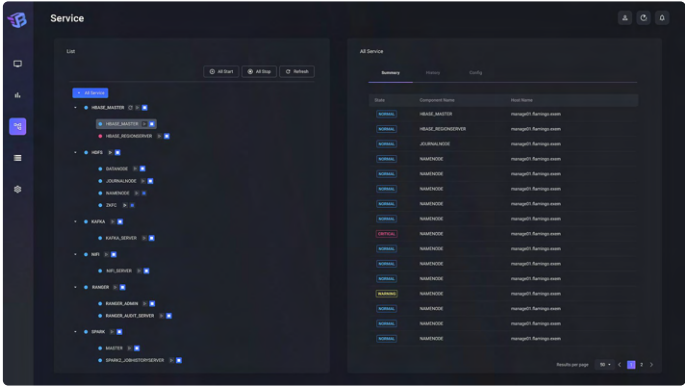
- Collect and visualize CPU, memory, disk, network usage and trends in real time
- Proactive response before availability degradation via threshold based alerts
- Instantly identify bottleneck points by comparing resource status across nodes/ services



3 Service Management

Easily achieve operations standardization and optimal configuration in a single screen

- Conveniently modify and deploy service settings via Web UI
- Derive optimal parameters by comparing execution change history
- Easily operate Apache Hadoop Ecosystem with multiple components via standard procedures



Architecture

1 Customer Environment

- Apache Hadoop 2.0, 3.0 based Ecosystem
- Manage Kafka, Hive, HDFS, Spark, HBase, Zeppelin
- Hadoop cluster hosts and service configurations

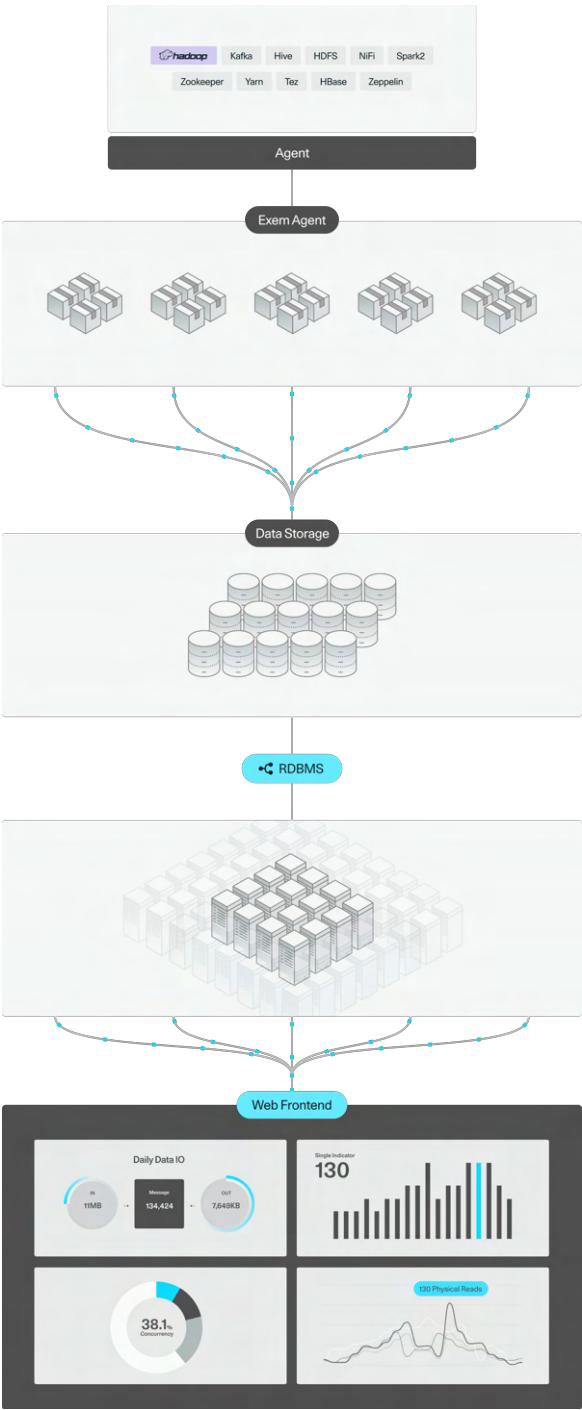
2 Data Collection & Processing

Collection & Monitoring

- Collect all Hadoop cluster metrics in real time
- Monitor core services: Resource Manager, NameNode, YARN
- Detect workflow errors and verify history

Management & Operations

- Modify and deploy service settings via Web UI
- Create/modify/delete HDFS directories and files (HDFS browsing)
- Provide Hive database management and query editor



3 Dashboard

- Visualize Hadoop cluster nodes via 3D dashboard
- Provide service and threshold status alerts
- Operate security policies through permission settings

Platform Specs

Web Browser

Chrome, Edge
Resolution: 1920×1080 (recommended) / 1440×900 (minimum)

Worker Server

OS: Linux (RedHat, CentOS, Ubuntu, etc.)
CPU: 32 Core (recommended)
RAM: 256GB (recommended)
DISK: 512GB SSD * 2 / 8TB HDD (recommended)
NODE: 4 units (recommended)

Master Server

OS: Linux (RedHat, CentOS, Ubuntu, etc.)
CPU: 16 Core (recommended)
RAM: 128GB (recommended)
DISK: 512GB SSD * 2 / 2TB HDD (recommended)
NODE: 2 units (recommended)



Why Woodpecker Stands Out

One Stop Analytics Platform
from AI Model Development to Deployment

Woodpecker is a Kubernetes based self service AI analytics platform. Start projects with just a few clicks using pre-configured development environments and diverse analytics tools including Jupyter and VS Code, dramatically simplifying model training and deployment. As a result, reduce operational burden with flexible resource allocation and focus on deriving insights for faster, more accurate outcomes.

Project Edit

Service Information

User	Woodpecker	Update Date	2025-12-19 10:08:59
Project Name *	Urban Bicycle Demand Prediction		
Project Description	Bike Sharing Demand Prediction Model based on Time, Weather, and Day of Week Data		
Image Name	woodpecker_v2.10 (python3.9)	Server	woodpecker2

Assign Resource

Resource	MAX	Assignment	Available
CPU	20	12	8
Memory	77.8	26	51.8
Disk	10246.6	355	9891.6

CPU

Minimum specifications
2 Core

0.02.0

MEMORY

Minimum specifications
4 GB

0.04.0

DISK

Minimum specifications
20 GB

20.0

Related Products

Use AI analytics platform Woodpecker together with Big Data operations platform EBIGs. Perform the entire workflow from data collection, loading, analysis, to visualization.



Product Highlights



Visualize from Key Metrics to Issue Identification

Support diverse environments including Jupyter, VSCode. From data patterns to workflow monitoring and issue identification.



AI Model Development to Deployment

Customize AI development with automated training scheduling, auto generated deployment APIs, and integrated development environments.



Monitor Deployed Model Performance in Real Time

Visualize key performance metrics (latency, throughput, success rate) for deployed AI models in real time dashboards.



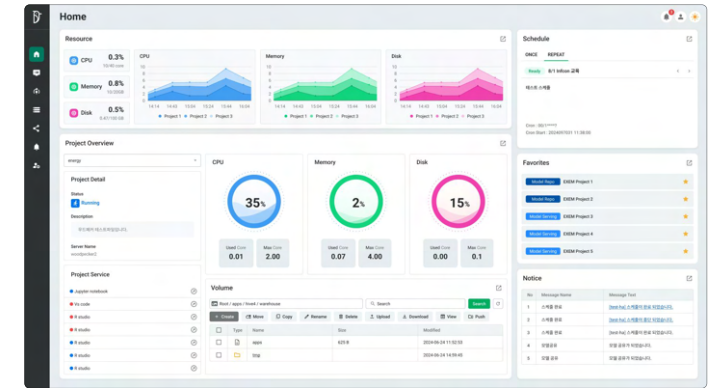
Data Protection with Sophisticated Surveillance

Prevent unauthorized access to sensitive information through user customized environment isolation and policy based access control.

1 Unified AI Model Development Environment Management

Integrated interface for real time project resource monitoring (CPU, Memory, Disk) with instant analytics tool access. View training schedules at a glance with model favorites and alert features.

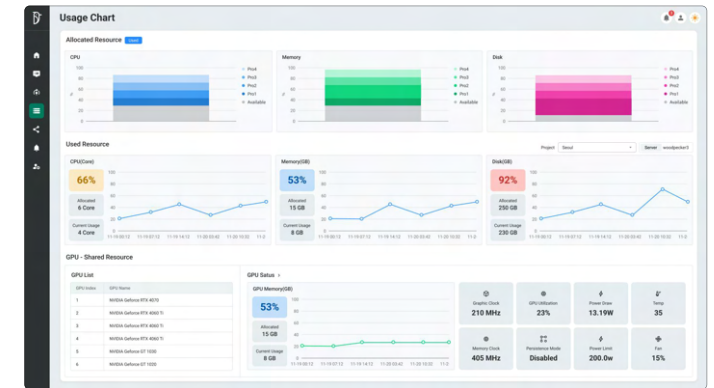
- Real time resource usage for projects/models
- Execute key project/model functions
- Manage model training schedules
- Provide user notifications



2 Resource Monitoring

Display AI development environment and deployment task resources (CPU, Memory, Disk) as intuitive graphs. Provide hourly resource usage and GPU details for efficient resource management.

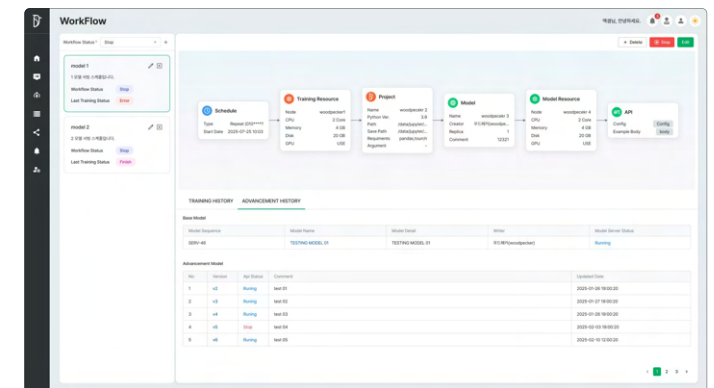
- Manage allocated resources
- Provide hourly resource usage by project/model
- Provide GPU usage and detailed information



3 Model Training Scheduling

View all user schedules intuitively. See reserved/recurring schedules chronologically with user, server, and schedule details.

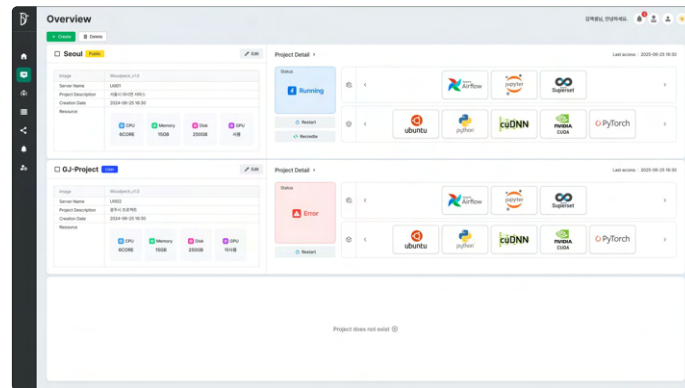
- Create model training schedules (reserved/recurring execution)
- View scheduling history and logs
- Provide all user schedule timeline



4 AI Model Development

Build AI model development environments easily with Project Overview. Verify project status in real time and perform data exploration, development, and training one stop.

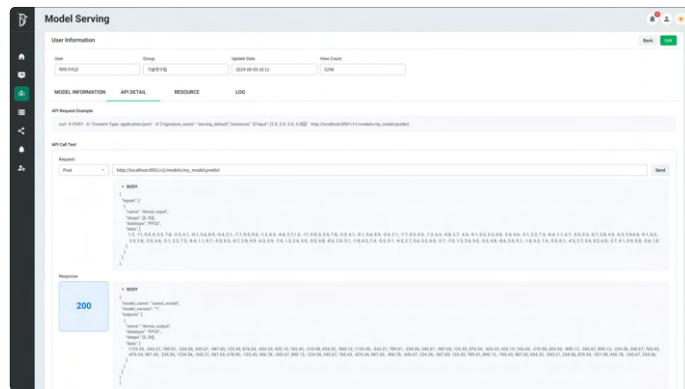
- Create AI model development environments with base images
- Configure and modify AI model development environment resources
- Provide environments for data exploration and model development/training through analytics tools



5 Model Serving

Attach models and enter basic info to auto generate API endpoints. Call models via provided URL and receive real time results instantly.

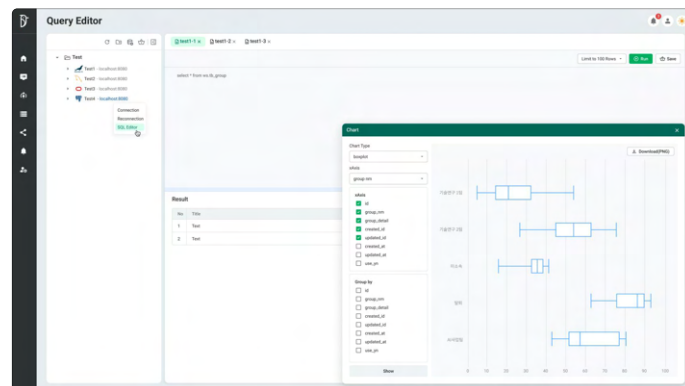
- Provide model deployment via auto generated APIs
- Provide model API call functionality
- Provide model execution logs



6 Data Query and Visualization

Access databases directly to execute SQL and visualize data in various formats. Easily analyze complex patterns and discover hidden anomalies.

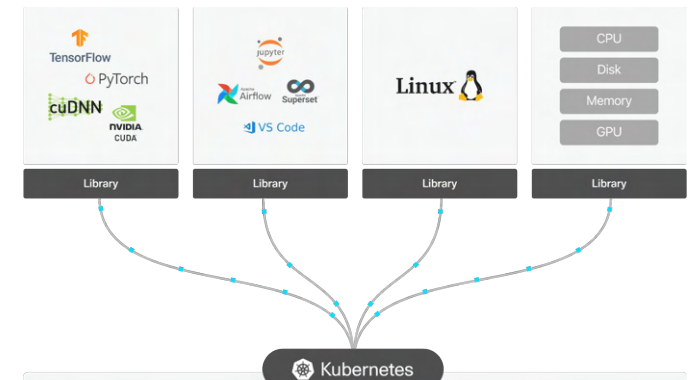
- Provide SQL functionality with database access
- Intuitive data analysis through data visualization
- Provide Bar/Line/Scatter/Heatmap/Boxplot charts



Architecture

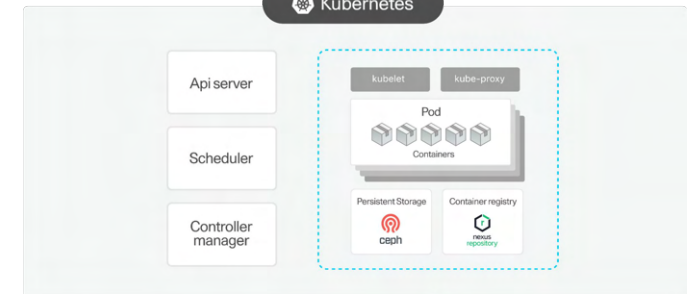
1 Analytics Environment Spec Configuration

- Pre installed major AI frameworks: TensorFlow, PyTorch, cuDNN
- Support diverse analytics tools: Jupyter, VS Code
- Flexible resource allocation: CPU, Memory, GPU



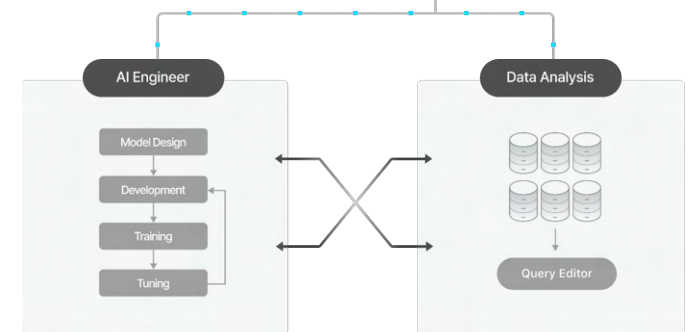
2 Large Scale Analytics Environment Build

- Kubernetes based clustering and container management
- Build AI model development environments easily without complex infrastructure setup
- Real time resource usage monitoring by project/ model



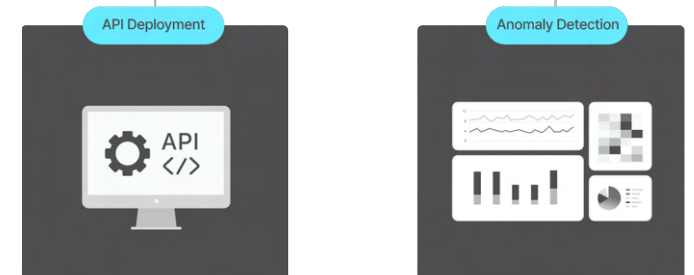
3 Modeling / Data Analysis

- One stop support from model design to development, training, and tuning
- Provide model training scheduling and auto execution
- DB access and SQL execution via Query Editor



4 Model Deployment / Visualization

- Easy model deployment via auto-generated APIs
- Support real time prediction and inference
- Provide data visualization: pattern analysis, anomaly detection



Chrome, Edge
Resolution: 1920×1080 (recommended) / 1440×900
(minimum)

Worker Server

OS: Linux (RedHat, CentOS, Ubuntu, etc.)
CPU: 128 Core (recommended)
RAM: 512GB (recommended)
DISK: 960GB SSD * 2 / 12TB HDD * 2 (recommended)
NODE: 4 units (recommended)

Master Server

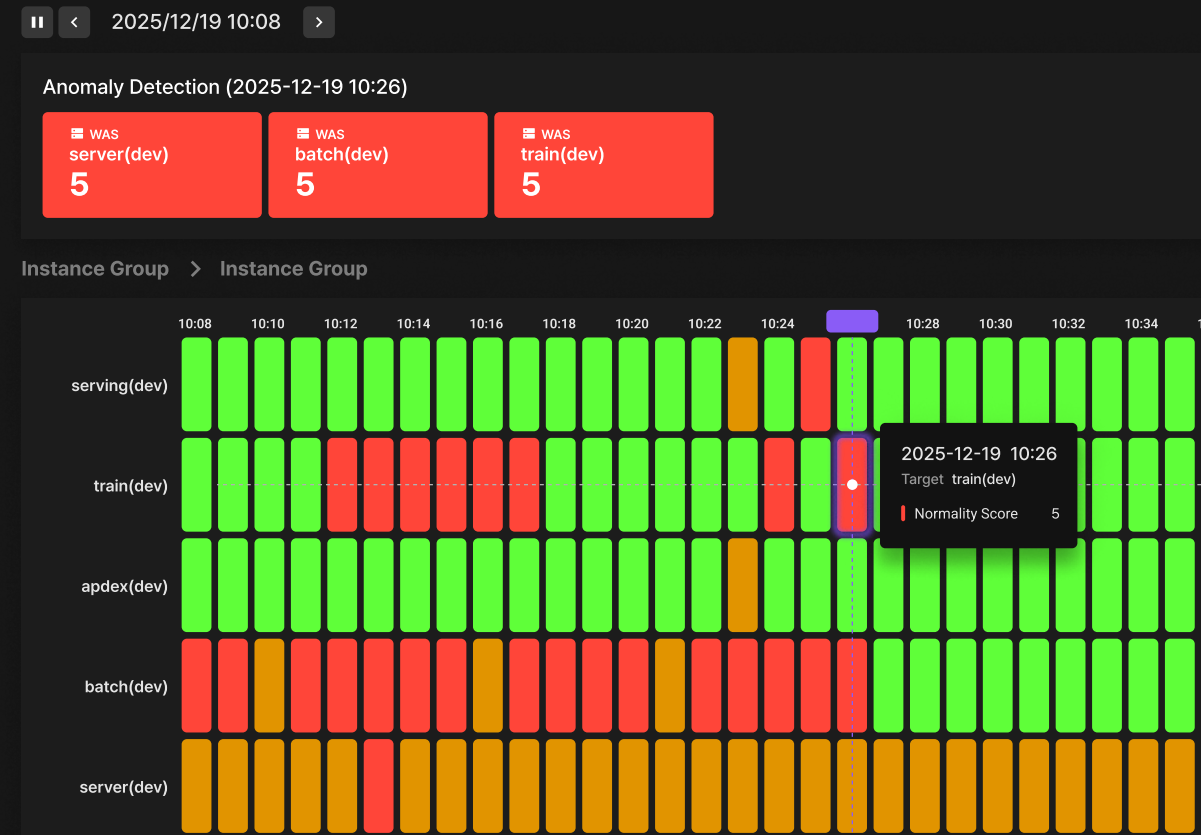
- OS: Linux (RedHat, CentOS, Ubuntu, etc.)
- CPU: 32 Core (recommended)
- RAM: 128GB (recommended)
- DISK: 960GB SSD * 2 / 12TB HDD * 2 (recommended)
- NODE: 3 units (recommended)



Why XAIOps Stands Out

The Future of Intelligent IT Operations: Real Time Anomaly Detection and Incident Prediction

XAIops is an AIOps solution that analyzes server, DB, network, and log data with AI to instantly detect anomalies and predict incidents. Rapidly identify root causes through event correlation analysis and runbook/playbook based standard response procedures. Reduce alert noise and repetitive tasks while shortening recovery time.



Customer Stories

Financial Institution

Proactive Financial Control with AI Prediction

Integrated collection, processing, and analysis of diverse IT operations data with an AIOps based solution. Achieved proactive incident response through ML based anomaly detection and real time decision making, improving operational efficiency through unified cross department root cause analysis and workforce reallocation.

Public Institution

Enterprise Integration for Fast, Accurate ICT Control

Deployed a platform integrating enterprise wide monitoring data to establish centralized performance and incident analysis. Enhanced ML/DL based anomaly detection and root cause analysis accuracy for rapid control, proactive response, and unified communication between development and operations.

Product Highlights



Capture Hard to Spot Anomalies in Real Time

AI models learn structured and unstructured data to detect performance anomaly patterns and analyze correlations. Identify abnormal logs with ease.



Analysis Features Showing Problem Patterns

Learn historical data to identify recurring problem patterns and causes. Predict situations 30–60 minutes ahead for preemptive action.



Diagnose Real Root Causes Quickly with AI

Diagnose root causes by referencing existing data and related metrics through real time intelligent detection. Identify issues faster and more accurately.



AI Chatbot Diagnosing Issues

Diagnose issues in real time with 'QURI', an LLM based chatbot. Get optimal answers by context: system status, incident diagnosis, load prediction.



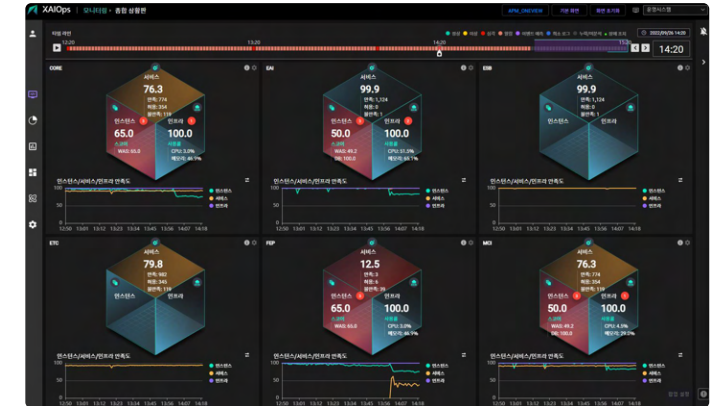
Real Time Monitoring

AI diagnoses entire IT system status in real time, detects anomalies, and alerts users. Create custom dashboards to view applications, infrastructure, databases, and more at a glance for easy management of complex systems.

1 Unified Status Board

Diagnose entire IT operations status at a glance with AI based satisfaction index

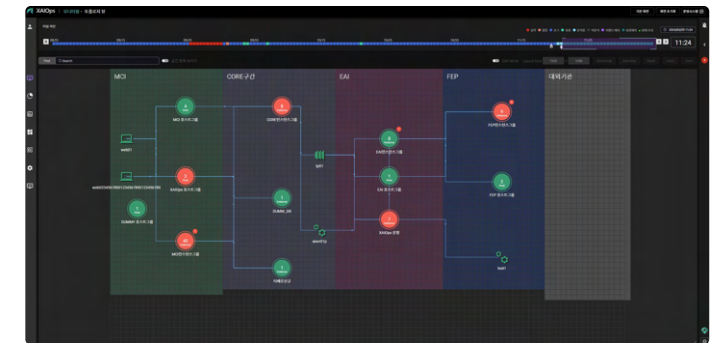
- Display combined status: applications, instances (WAS, DB), infrastructure (CPU, Memory)
- Provide 'Satisfaction Index' calculated from Apdex, TPS, resource usage
- Enable rapid response with real time alerts and auto anomaly detection



2 Topology View

View call paths, status, and anomaly timing across all segments in one screen

- Visualize relationships and data flows by domain (Web/WAS/EAI-ESB/FEP/DB)
- Distinguish host and instance group status, alert history, anomaly detection
- Enhance recognition efficiency by visualizing anomaly levels with colors and icons



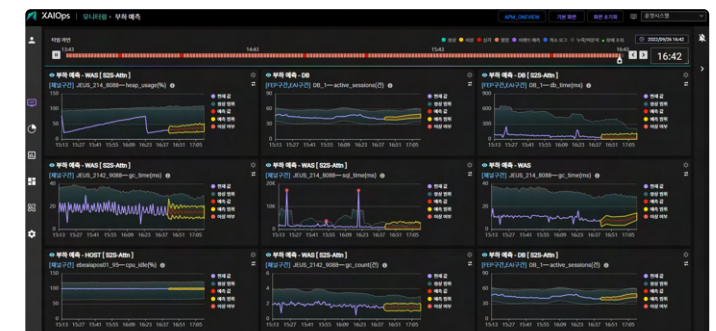
Load Prediction and Anomaly Detection

AI learns IT system load and status to predict future issues and detect anomalies deviating from normal ranges in real time for rapid response. Auto-analyze key metrics (server, application, database, network) to alert risk signs early.

1 Load Prediction

Predict future traffic and resource demand through historical learning to prevent overload

- Predict surge timing via WAS metrics: transaction volume, response time, error count
- Pre-identify peak periods through DB pattern learning: processing count, Lock sessions
- Plan capacity expansion and distribution via usage prediction: CPU, memory, disk, network



2 Anomaly Detection

Instantly detect and alert changes deviating from AI defined normal ranges

- Monitor core metrics in real time: WAS, DB, infrastructure
- Auto learn and apply normal range (Base Line)
- Immediately recognize response delays, error spikes, excessive resource usage



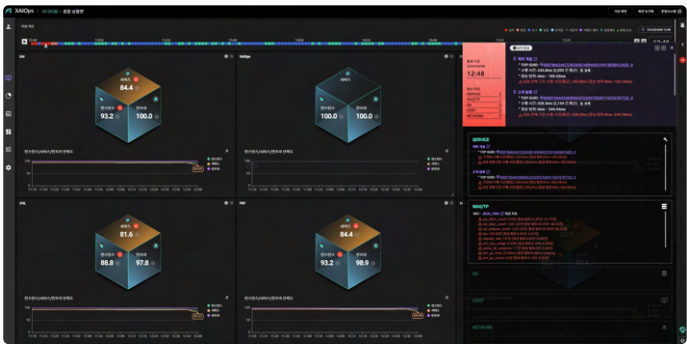
Incident Root Cause Analysis

AI detects incident points in real time during service anomalies and analyzes correlations and flows across metrics and logs to quickly identify root causes. Auto collect and analyze data across all domains (WAS, DB, network, unstructured logs) to dramatically reduce incident response time.

1 Incident Root Cause Analysis

Identify incident points and causes through correlation analysis to shorten recovery time

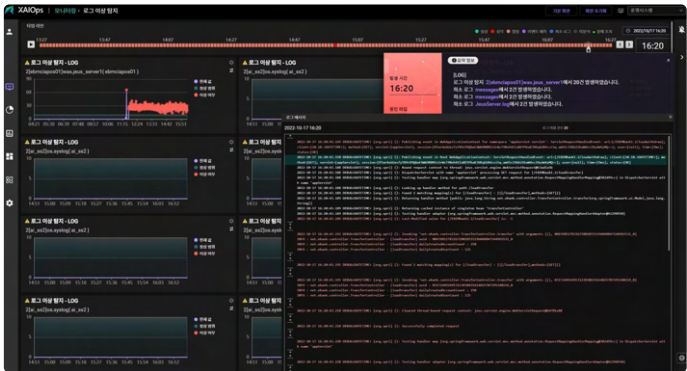
- Identify problem points through precision analysis of segments/items at occurrence time
- Infer inter system impact relationships based on metric change flows
- Support rapid resolution by suggesting priority action points



2 Causal / Correlation Analysis

Find correlations across metrics and logs at the same time to identify 'cause→effect' flows

- Identify impact scope through cross analysis: WAS, DB, network, logs
- Auto extract and visualize highly correlated items
- Use propagation path analysis for recurrence prevention



3 Log Anomaly Detection Analysis

Learn unstructured logs to capture anomaly patterns in real time and correlate with metrics

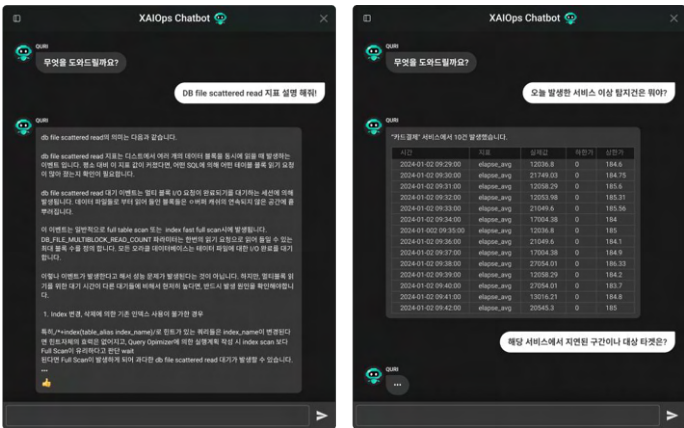
- Auto identify anomaly signs across diverse logs (Biz, WAS, System, etc.)
- Provide root cause clues based on error codes and recurring patterns
- Comprehensive diagnosis by linking related metrics and events



4 LLM based Chatbot (QURI)

Operations specialized AI chatbot for natural language queries and instant screen/function access

- Rapidly query alarm/incident/performance history by period and target
- Instantly navigate to requested dashboard, service, or instance screens
- Provide detailed answers on features, settings, and alert policies



Architecture

1 Customer Environment

- Collect data across entire IT infrastructure
- Unified collection of structured data (performance metrics) and unstructured data (logs)
- Support integration with Exem monitoring solutions

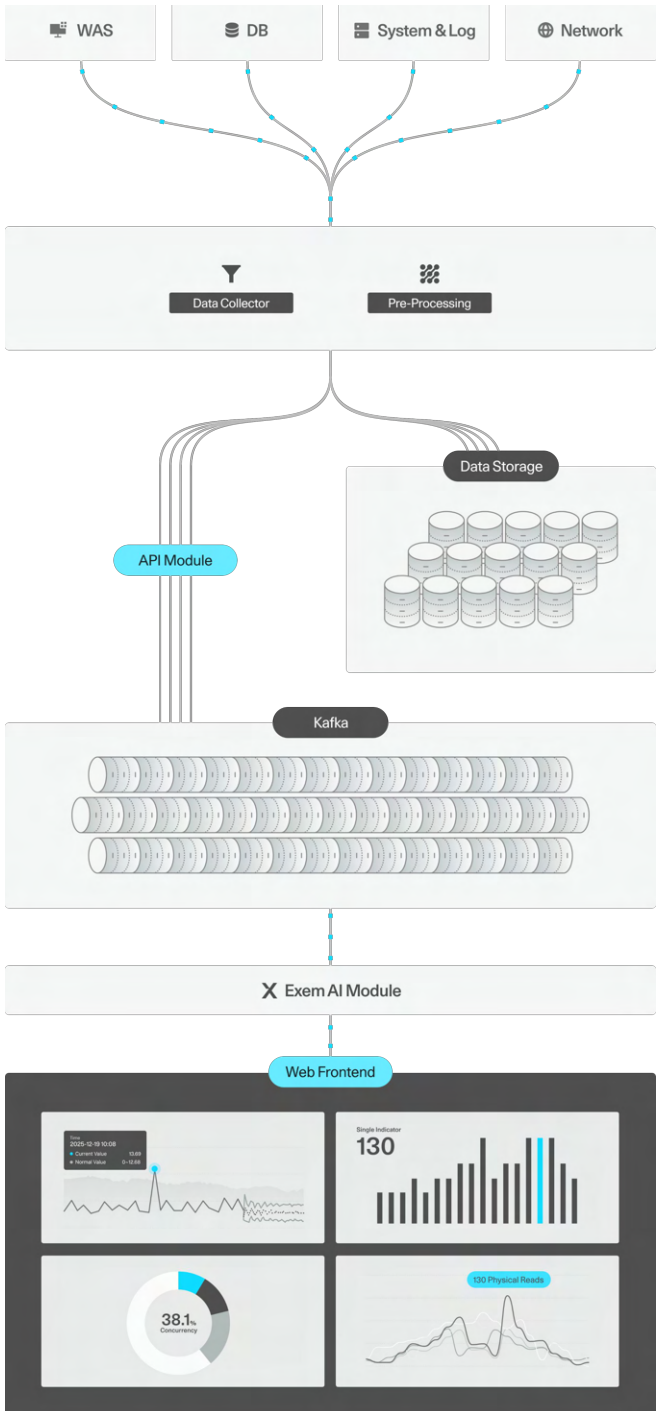
2 Data Collection & Processing

Collection & Preprocessing

- Collect real time performance metrics via Data Collector
- Refine data for AI learning at Pre Processing stage
- Process large scale streaming data via Kafka

AI Analysis

- Auto generate trust intervals (Base-line) based on deep/machine learning
- Real time anomaly detection and load prediction through historical data learning
- Provide incident root cause through causal/correlation analysis



3 Dashboard

- Provide unified status board for service, instance, and infrastructure perspectives
- Smart Alert adapting to dynamic load conditions based on AI learning
- Support natural language queries via LLM based chatbot 'QURI'

Platform Specs

Web Browser

Chrome 75 or higher

Data Collection / Processing Server

Based on 1 Unit

OS: Linux Kernel 4.18 or higher (64bit)
CPU: 24Core (recommended) / 12Core (minimum)
RAM: 256GB (recommended) / 128GB (minimum)
HDD: 8TB (recommended) / 4TB (minimum) - SSD (recommended)
Disk capacity may vary based on customer data collection and storage volume

AI Training Server

Based on 1 Unit

OS: Linux Kernel 4.18 or higher
CPU: 24Core (recommended) / 12Core (minimum)
GPU: NVIDIA H100 *1ea - GPU quantity may vary based on customer data training volume
RAM: 256GB (recommended) / 128GB (minimum)
HDD: 8TB (recommended) / 4TB (minimum), SSD (recommended)
Disk capacity may vary based on customer data collection and storage volume

Data Everywhere,
Make it Matter