



Security Fact Sheet

Edwin AI is LogicMonitor’s AI-powered agent designed to enhance ITOps workflows through streamlining incident automation with event intelligence and accelerating incident resolution with AI-powered insights. This document provides an overview of LogicMonitor’s approach to securing and governing its AI-powered capabilities, including Edwin AI. LogicMonitor may update the most current Edwin AI information from time to time, so for the most current information, please consult the [LogicMonitor Trust Center](#), your account executive, or our [support team](#).

Date	Version	Notes
January 21, 2025	1.0	Initial release
March 25, 2025	1.1	Expanded coverage on data security, control and access as well as model integrity and trust.
May 21, 2025	1.2	Minor updates for clearer language.
May 29, 2025	1.3	Minor updates for clarity related to AI Agent release.
July 10, 2025	1.4	Details added on additional model providers.
July 18, 2025	1.5	Updated details for APAC customers to accurately reflect models used in Edwin AI prompts.
August 8, 2025	1.6	Added specific questions around storing user login and third party platform credentials.
September 10, 2025	1.7	Added details about compliance with EU AI Act .
February 15, 2026	1.8	Added details about Edwin AI’s high availability architecture.
February 27, 2026	1.9	Added details about auditability and governance.

LLM Selection Process

WHAT LARGE LANGUAGE MODEL (LLM) DOES EDWIN AI UTILIZE CURRENTLY?

LogicMonitor Edwin AI is an AI-powered agent designed to enhance ITOps workflows using Large Language Models (LLMs) and retrieval-augmented generation (RAG). LogicMonitor may employ a variety of models that have undergone compliance reviews. Currently, Edwin AI Agent features use OpenAI models directly and Anthropic’s Claude models hosted via Amazon Bedrock. The Edwin AI platform is designed to be LLM agnostic and, depending on the specific use case and fit, is capable of integrating with different LLMs.

IS EDWIN AI DEPENDENT ON A CERTAIN LLM PROVIDER?

LogicMonitor’s Edwin AI utilizes an agentic architecture that is generally LLM agnostic and can, in principle, plug into many LLM providers. Edwin AI may use different LLMs based on the task that needs to be completed. As such, these models may change as we evaluate which models will best fit the use cases.

CAN I CHOOSE WHICH LLM TO UTILIZE?

We do not currently offer customers the ability to choose which models they would like to interact with. Customers can choose a hosting provider, if they have specific data restrictions. The Edwin AI team evaluates various models to determine the best fit for each workflow. However, if you have a specific LLM requirement, please share this feedback with our product and account team.

Data Security, Control & Access

WHAT INFRASTRUCTURE IS USED TO HOST THE MODELS?

Edwin AI's agentic features are available on OpenAI and Amazon Bedrock. If there are hosting provider restrictions, please let your account team know so the appropriate hosting provider can be selected.

Our proprietary Machine Learning Models used for Event Correlation are deployed within our AWS environment. Users are able to create and adjust the Correlation Models from the Edwin AI user interface.

IS MY DATA USED TO TRAIN AI MODELS?

As an administrative control, LogicMonitor maintains agreements with applicable LLM providers that prohibit customer data from being used to train third-party LLMs, such as OpenAI's or Anthropic's models. Your data will also not be used to train or improve AWS's products or any other third-party services. Your data will not be used to cross-train other models used by other customers.

IS MY DATA SHARED WITH ANY THIRD PARTIES?

No, except for internet search (see [What about internet search?](#)) and our approved sub processors (including the applicable LLM) that are necessary to provide the service (see [LogicMonitor's Data Handling Supplement](#)), your data will not be shared outside of LogicMonitor or between other customers.

User queries with the Edwin AI Agent and their associated context (relevant insight and alert metadata) are passed to and processed by OpenAI. For OpenAI-based workflows, LogicMonitor has enabled Zero Data Retention (ZDR). As a result, data processed through Edwin AI is not retained by OpenAI and is not used to train OpenAI models. Read more about Open AI policies [here](#).

Any data processed in Amazon Bedrock is encrypted and stored at rest in the specific AWS Region. Data is not shared with any model providers. Amazon is committed to the responsible use of AI and has an automated abuse detection mechanism as outlined by their [Acceptable Use Policy](#) and [Responsible AI Policy](#) but is not used for training purposes.

WHAT ABOUT INTERNET SEARCH?

Edwin AI may use an internet search tool to respond to certain user queries. When this occurs, the content of the customer's question is sent to an external search provider to retrieve relevant information.

IS MY DATA SHARED WITH OTHER CUSTOMERS?

No, your data is not shared with other customers. Customer data is split into different datastores that do not overlap with other customers.

IS MY DATA COMBINED WITH OTHER CUSTOMERS' DATA FOR EVENT CORRELATION?

Edwin AI ensures complete data isolation for each customer. There is no co-mingling of data. This applies to Managed Service Providers (MSPs) and their tenants as well. Each tenant's data remains logically separated, which is designed to ensure that correlation does not happen across customers.

IS DATA ENCRYPTED AT REST?

Edwin AI encrypts data in transit and at rest (HTTPS TLS 1.3) between LogicMonitor Envision and ServiceNow.

HOW ARE CREDENTIALS STORED WITHIN EDWIN AI?

Customer credentials for third-party platforms (such as ServiceNow) are securely stored in AWS Secrets Manager using industry best practices. These secrets are stored in the same AWS region as the customer's Edwin AI portal environment. User login credentials are never stored separately in Edwin AI; authentication is handled entirely through the LogicMonitor platform via SSO. LogicMonitor API tokens used for data ingestion (such as events and enrichment properties) are also stored securely in AWS Secrets Manager in the same region as the portal.

Data Retention

WHAT DATA IS STORED IN EDWIN AI?

A range of data is stored within Edwin AI such as first-party data from LogicMonitor like events, alerts, and insights along with external data such as ServiceNow incidents, knowledge base articles from third-party sources, and other data that the customer enables to be sent to Edwin AI. Customers can limit or stop sending data at any time. LogicMonitor asks that all Edwin AI users follow the [AI Acceptable Use policy](#).

Edwin AI utilizes Langfuse for observability and trace-level logging of AI Agent conversations and interactions. Logged data includes metadata such as prompt text, completion text, insight context, internet searches from the web search tool, latency and usage metrics. This data is used exclusively for debugging, quality assurance, and performance optimization. All data is stored and processed in region.

Langfuse is deployed using a region-specific architecture:

- **AMER and EMEA:** Edwin AI uses the Langfuse Cloud service, which provides regional data residency and ensures all data is stored and processed within the respective geographic region.
- **APAC:** Edwin AI continues to use a self-hosted Langfuse deployment, ensuring all logged data remains within LogicMonitor's controlled infrastructure in the APAC region.

In all cases, data is processed in compliance with LogicMonitor's security and privacy policies. No logged data is shared with third parties or used by Langfuse for model training.

Langfuse's built-in data retention controls allow us to define how long this observability data is stored. We retain only the data necessary for operational and security purposes, with a configurable minimum retention window of three days. You can learn more about Langfuse's retention capabilities [here](#).

User queries with the Edwin AI Agent and their associated context (relevant insight and alert metadata) are passed to and processed by OpenAI. For OpenAI-based workflows, LogicMonitor has enabled Zero Data Retention (ZDR). As a result, data processed through Edwin AI is not retained by OpenAI and is not used to train OpenAI models. Read more about Open AI policies [here](#).

HOW DO LLM PROVIDERS HANDLE MY DATA?

OpenAI: When using OpenAI directly, prompts and completions may be retained for up to 30 days strictly for abuse monitoring and safety. No data is used for training or shared with external vendors. See [OpenAI usage and retention policy](#) for more info.

Amazon Bedrock: Customers control their data and no data is stored in the Amazon Bedrock service. Prompts and completions are not stored or logged by Amazon Bedrock. See [Amazon Bedrock security](#) for more info.

HOW IS MY DATA STORED IN EDWIN AI?

Customer data is securely isolated into dedicated physical or logical datastores, ensuring each customer's data remains separate.

Data Processing & Compliance

WHAT DATA IS CURRENTLY PROCESSED BY EDWIN AI?

LogicMonitor's Edwin AI currently utilizes any data that is available within Edwin AI, which primarily includes alert information, enrichment information that is intended to better classify ingested information, user content that is sent via the AI Agent, and any additional information that has been configured to be sent to Edwin AI. LogicMonitor maintains contractual obligations with all approved subprocessors, conducts privacy impact assessment(s) where applicable, and protects personal data.

DOES THE EDWIN AI AGENT SUPPORT DATA HOSTING IN A PARTICULAR REGION?

Yes. Edwin AI is designed with data residency compliance in mind.

Customer data is processed and stored within regional infrastructure to comply with local data residency requirements. For OpenAI-based deployments, regional processing is aligned with OpenAI's data residency policies (see OpenAI's [documentation](#)).

If you select Amazon Bedrock as your underlying LLM platform, your content is encrypted and stored at rest in the AWS Region where the Amazon Bedrock service is deployed.

WHERE ARE THE USERS WHO WILL ACCESS OUR SYSTEMS AND DATA GEOGRAPHICALLY LOCATED?

LogicMonitor operates globally. Your system and data will only be accessed when necessary by authorized product, customer support, or other technical team members. If you have specific geographic restrictions on data access, please inform us, and we will make every effort to accommodate your requirements. Further details on our data handling can be found [here](#).

HOW DOES EDWIN AI HANDLE PERSONALLY IDENTIFIABLE INFORMATION (PII) AND SENSITIVE DATA?

Edwin AI is designed to minimize the use of personally identifiable information (PII) and sensitive data. By default, Edwin AI only processes PII that is necessary for understanding and resolving requests. This may include incident fields that reference names, email addresses, phone numbers, roles, and usernames.

Customers are responsible for managing the data they send to Edwin AI. If certain fields contain PII or other sensitive data that customers prefer not to share, those fields can be excluded from the data sent through integrations.

DOES EDWIN AI COMPLY WITH AI LAWS?

Yes. Edwin AI is not intended as a safety system component, nor is it designed or used to make decisions about an individual (such as use in employment).

While the [EU AI Act](#), and similar U.S. legislation is still in the process of being implemented, we have aligned our AI use and development practices with its requirements and underlying principles of safety, accountability, and transparency. To the best of our knowledge, LogicMonitor's Edwin AI service does not fall into a prohibited use case, and is not designed to be used as a safety component in any common market products that would require a conformity assessment.

- **No model training or tuning:** We do not train or fine-tune AI models.
- **No use of personal data for training:** We do not allow customer data or personal data (as a data controller) to be used for training by any LLM.
- **AI risk and compliance reviews:** We conduct Privacy Impact Assessments (PIAs) and AI compliance assessments for all new AI use cases to ensure responsible adoption.
- **Ban of high-risk and prohibited use cases:** We do not onboard or allow the use of AI tools for high-risk or prohibited scenarios under the EU AI Act, including internal experimentation.
- **Product use and risk classification:** To the best of our knowledge, our product is not used in any high-risk or prohibited activities as defined by the EU AI Act.

In short, our approach ensures that our use of AI technologies is compliant with the EU AI Act's framework and consistent with protecting individual rights and organizational accountability.

AI Model Integrity & Trust

WHAT SAFEGUARDS ARE IN PLACE TO PREVENT AI-GENERATED MISINFORMATION OR HALLUCINATIONS?

Edwin AI incorporates several safeguards that are designed to reduce the risk of misinformation or “hallucinated” responses. Retrieval-augmented generation (RAG) is utilized to fetch information in real-time from trusted sources, which is a measure intended to ensure the outputs are grounded in factual data. The retrieved information is accompanied by reference source links, allowing users to verify the origin and credibility of content. Additionally, Edwin AI employs a rigorous evaluation process ensuring responses are continuously evaluated to maintain reliability and factual correctness (subject to the accuracy of the customer’s data).

HOW DOES EDWIN AI MONITOR AI OUTPUTS FOR BIAS?

Edwin AI monitors AI outputs for bias as part of its production evaluation pipeline. AI investigations are evaluated using an LLM-as-a-Judge framework that assesses responses for bias as part of an overall quality evaluation. If potential bias is identified, the investigation’s evaluation score is reduced, and lower-scoring investigations are flagged for human review. Edwin AI also uses retrieval augmented generation (RAG) to ground responses in trusted enterprise data and employs safety guardrails to reduce the risk of harmful or unsupported content.

WHAT SECURITY MEASURES AND THREAT MITIGATION STRATEGIES ARE IN PLACE TO PROTECT DATA AND SERVICES?

Edwin AI employs multiple layers of security controls to protect customer data and ensure service integrity.

Key technical and organizational measures include:

- Cloud provider security: Edwin AI leverages the native security features of its underlying cloud infrastructure, including network segmentation, access controls, and continuous monitoring
- Penetration testing: annual third-party penetration tests are conducted to proactively identify and remediate potential vulnerabilities
- Vulnerability scanning
- Role-based access controls (RBAC) and Identity and Access Management (IAM)
- Availability control measures

AI Governance & Auditability

WHAT TYPE OF AUDIT TRAIL APPLIES TO MY DATA?

All Edwin AI Agent decisions are documented for internal-use only. All inputs and outputs include source references and citations from retrieved content, enabling traceability. Edwin AI utilizes Langfuse for observability and trace-level logging of AI Agent conversations and interactions. Logged data includes metadata such as prompt text, completion text, insight context, internet searches from the web search tool, latency and usage metrics. This data is used exclusively for debugging, quality assurance, and performance optimization.

DOES YOUR AI SYSTEM SUPPORT HUMAN-IN-THE-LOOP WORKFLOWS FOR REVIEWING AND APPROVING AI OUTPUTS?

All information generated by Edwin AI is presented as a recommendation. Humans retain full control over decision-making and can review, validate, and determine whether to implement any suggested actions.

HAS THE AI SYSTEM IMPLEMENTED MECHANISMS TO LIMIT THE RATE AND SCOPE OF AI AGENT ACTIONS?

Yes. Edwin AI Agents are limited to a defined set of tools developed and controlled by the LogicMonitor team. Many tools are opt-in, requiring explicit customer authorization and data provision (e.g., incidents, change requests, knowledge base articles). This ensures customers maintain control over which agent capabilities are enabled and the scope of actions permitted. Additionally, the EdwinAI Agent has strict guardrails in place to ensure content discussed is limited to the ITOps domain.

HOW ARE MODEL BEHAVIOR AND UPDATES GOVERNED?

Edwin AI integrates with market leading LLM vendors like OpenAI and Amazon Bedrock and therefore updates are not tracked within Edwin AI directly. We continuously monitor vendor roadmaps for model releases and new capabilities to ensure alignment with performance, security, and reliability standards.

All candidate model upgrades undergo structured offline evaluations against a defined benchmark suite (prompt sets, task-specific scenarios, and edge cases). Outputs are assessed for quality, determinism, latency, safety compliance, and regression risk. No model or prompt change is promoted without passing predefined acceptance criteria.

We follow a staged roll out process. Model upgrades, prompt modifications, and feature changes are first deployed to a Sandbox environment where customers can validate behavior and provide feedback prior to production release. Promotion to production occurs only after validation thresholds are met.

High Availability & Resilience Architecture

LogicMonitor's Edwin AI platform is specifically designed to handle high-volume, multi-portal alert environments with exceptional reliability and scalability. Our architecture can confidently manage your alert volumes through several key capabilities.

ADVANCED ALERT VOLUME MANAGEMENT

Edwin AI employs advanced deduplication technology that reduces alert noise by 90–95% efficiency by transforming repeated events into single, contextualized alerts. The platform can compress up to 95% of incoming alerts which dramatically reduces processing overhead while maintaining complete visibility of critical issues.

PROVEN PERFORMANCE WITH LARGE-SCALE DEPLOYMENTS

Our platform has demonstrated success handling environments with 150 million+ events per month, resulting in a 75% reduction in ticket volume through intelligent correlation and grouping.

ENTERPRISE-GRADE ARCHITECTURE

The platform features a cloud-native, distributed architecture with horizontal scaling capabilities to accommodate growing data volumes without performance degradation.

INTELLIGENT WORKLOAD DISTRIBUTION

The platform automatically balances processing loads across system components, preventing bottlenecks during peak alert periods while maintaining consistent performance. Edwin AI enables you to achieve exponential scaling of your IT operations without increasing headcount.

ROBUST HIGH-AVAILABILITY FEATURES

Our architecture includes built-in redundancy and failover mechanisms to ensure continuous operation even during maintenance or unexpected events, guaranteeing reliability for critical monitoring requirements.

Our current clients have successfully implemented Edwin AI across similar multi-portal environments with significant reliability improvements.