

# Fujitsu Kozuchi Enterprise AI Factory Introductory Materials

Fujitsu Limited

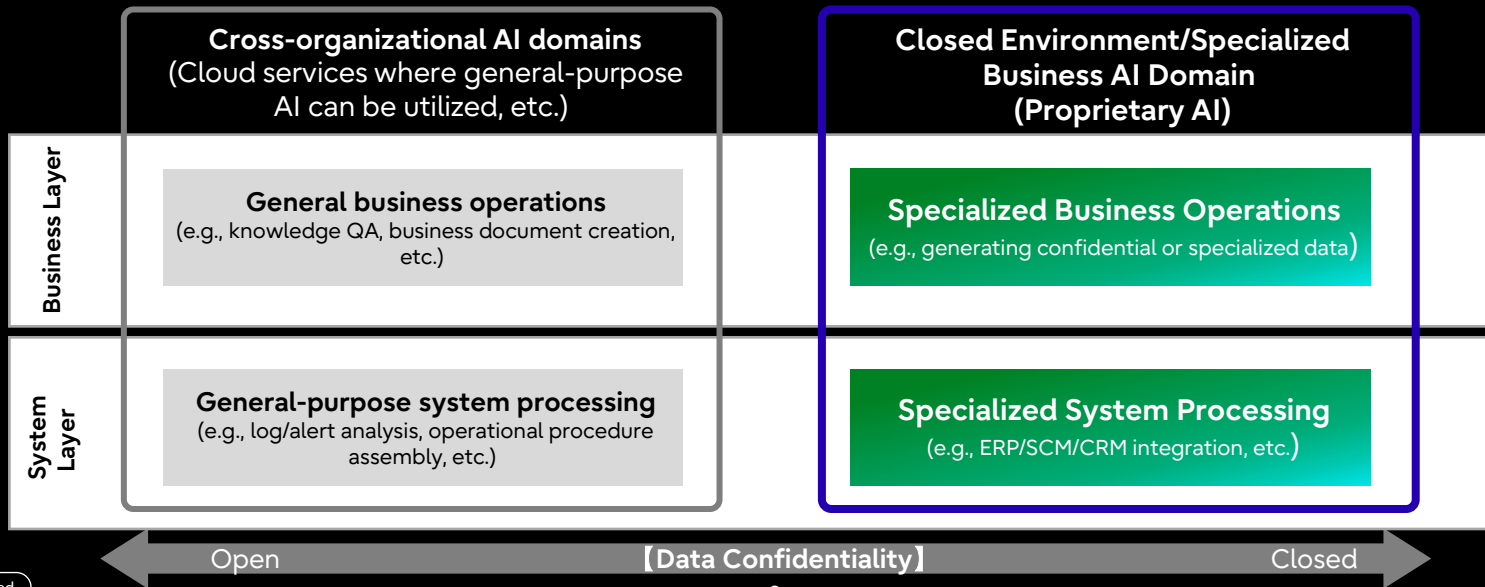
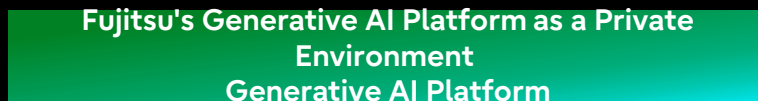
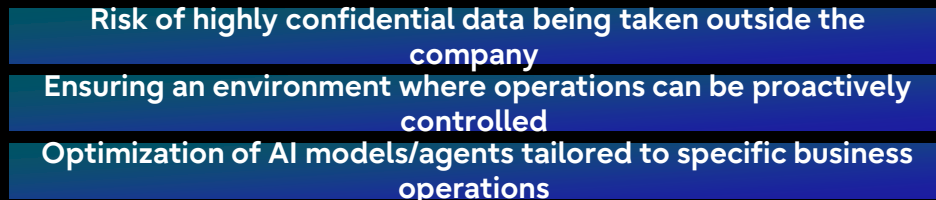


# Market Trends Surrounding Generative AI

# Fujitsu's Approach

uvance

- Configuring a generative AI platform as a private environment capable of securely handling "highly confidential data" presents a new challenge

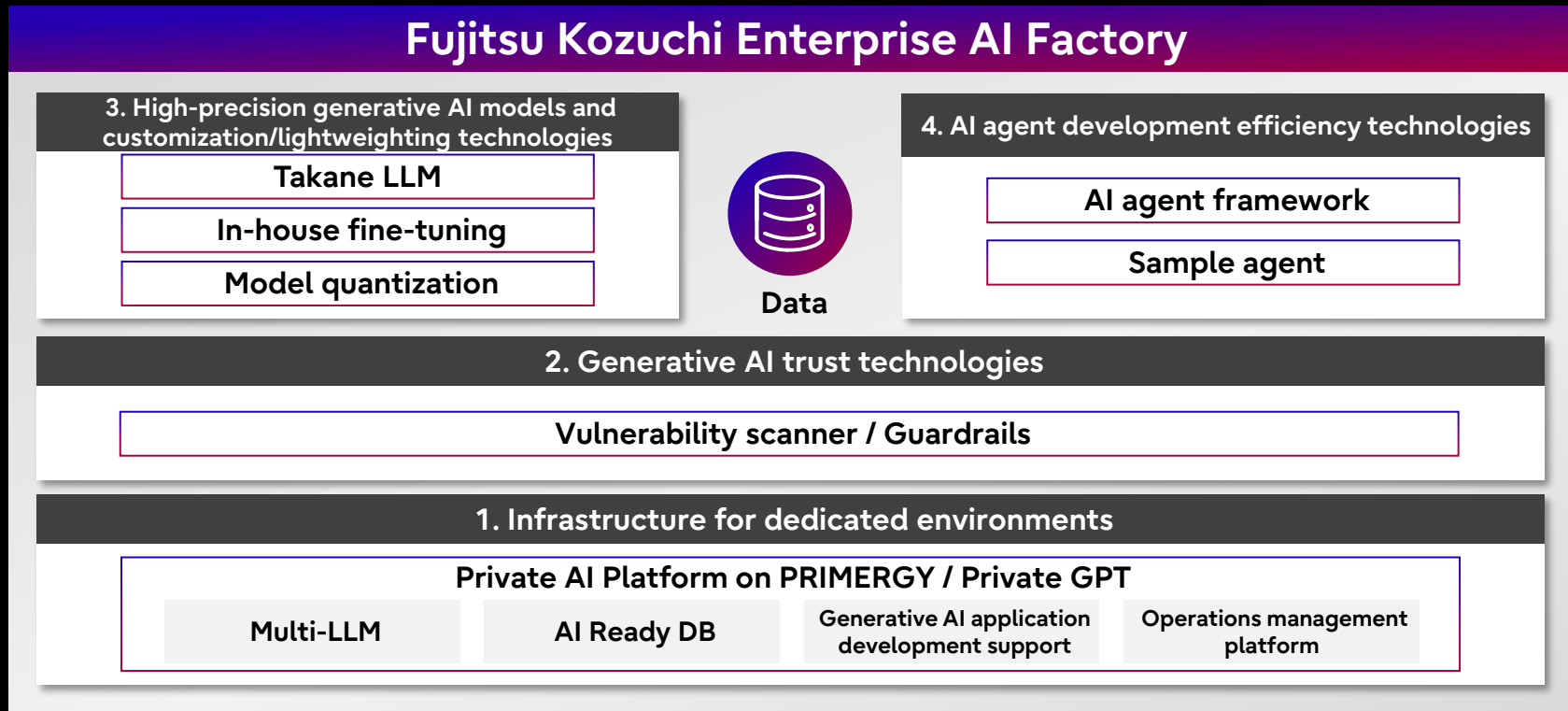


# Fujitsu Kozuchi Enterprise AI Factory Overview

# Fujitsu Kozuchi Enterprise AI Factory

uvance

- A proprietary AI platform that can be utilized on-premises as a means to achieve highly sovereign AI



- A proprietary AI platform deployable on-premises as a means to achieve highly sovereign AI.



## 1. Infrastructure for dedicated environments

- Achieving on-premise AI utilization that keeps data internal.
- Providing end-to-end support from introduction to advanced operation.
- Adopting PRIMERGY-based high-performance GPU infrastructure. \*



## 2. Generative AI trust technologies

- LLM vulnerability scanner supporting over 9,000 types of vulnerabilities.
- Guardrails to prevent malicious input and unsuitable output.
- Automated operation ensuring safety even for non-experts.



## 3. High-precision generative AI models and customization/lightweighting technologies

- Adopting the high-precision Japanese LLM "Takane" as the core.
- Developing and refining domain-specific models with in-house fine-tuning.
- Achieving low-cost and highly efficient AI operation through LLM quantization technology.



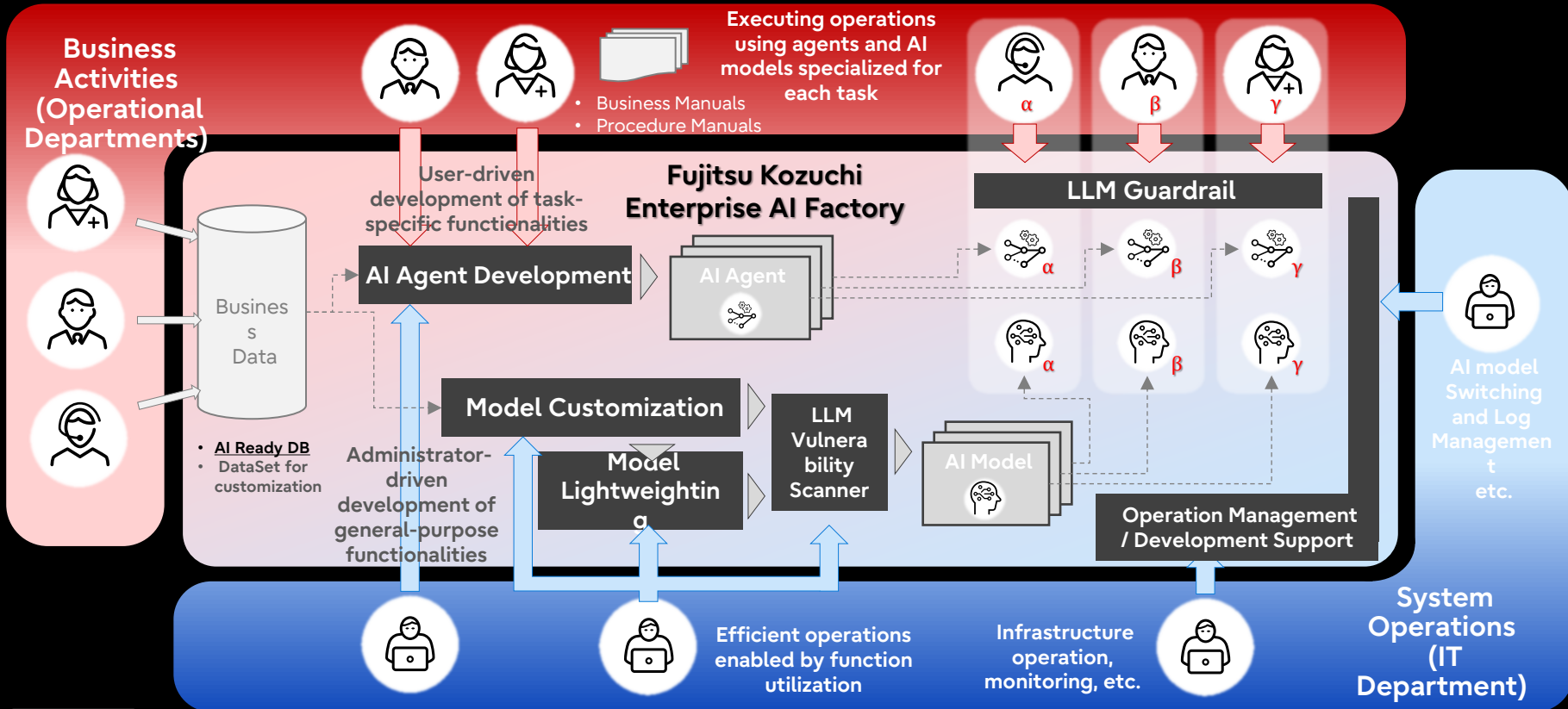
## 4. AI agent development efficiency technologies

- Quickly developing enterprise AI with low-code/no-code platforms. Rapidly build business AI
- Flexible integration with existing systems through MCP compatibility.
- Advanced utilization through collaborative actions of multiple AI agents.

\*Japan: Private AI Platform on PRIMERGY, Europe: Private GPT

# In-house production of specialized models driving business growth

uvance



# Key Target Areas

uvance

- Our primary targets are set to four domains: Manufacturing, Healthcare, Finance & Insurance, and Social Infrastructure & Public.
- A proprietary AI platform that enables the development, operation, and continuous advancement of AI tailored to industry requirements.

## Manufacturing



- Cross-handling drawings, manuals, and equipment data generated in each stage of design, production, and maintenance, and iteratively deploying them to on-site operations.
- Accumulating and reusing knowledge related to design intent understanding, quality analysis, robotics, and automated equipment across all processes.

## Healthcare



- Securely importing medical records, examination reports, and internal documents, and continuously utilizing them throughout the entire business process.
- Streamlining and enhancing healthcare operations through the standardization and reuse of document generation, search, and summarization, based on on-premise and closed environments.

## Finance and Insurance



- Centralized management of business data such as policies, regulations, review materials, and inquiry histories, embedding them into business processes for continuous utilization.
- AI utilization predicated on stable operation and scalability, achieved through the standardization of inquiry responses and document verification, and the reflection and updates of business rules.

## Social Infrastructure & Public Services



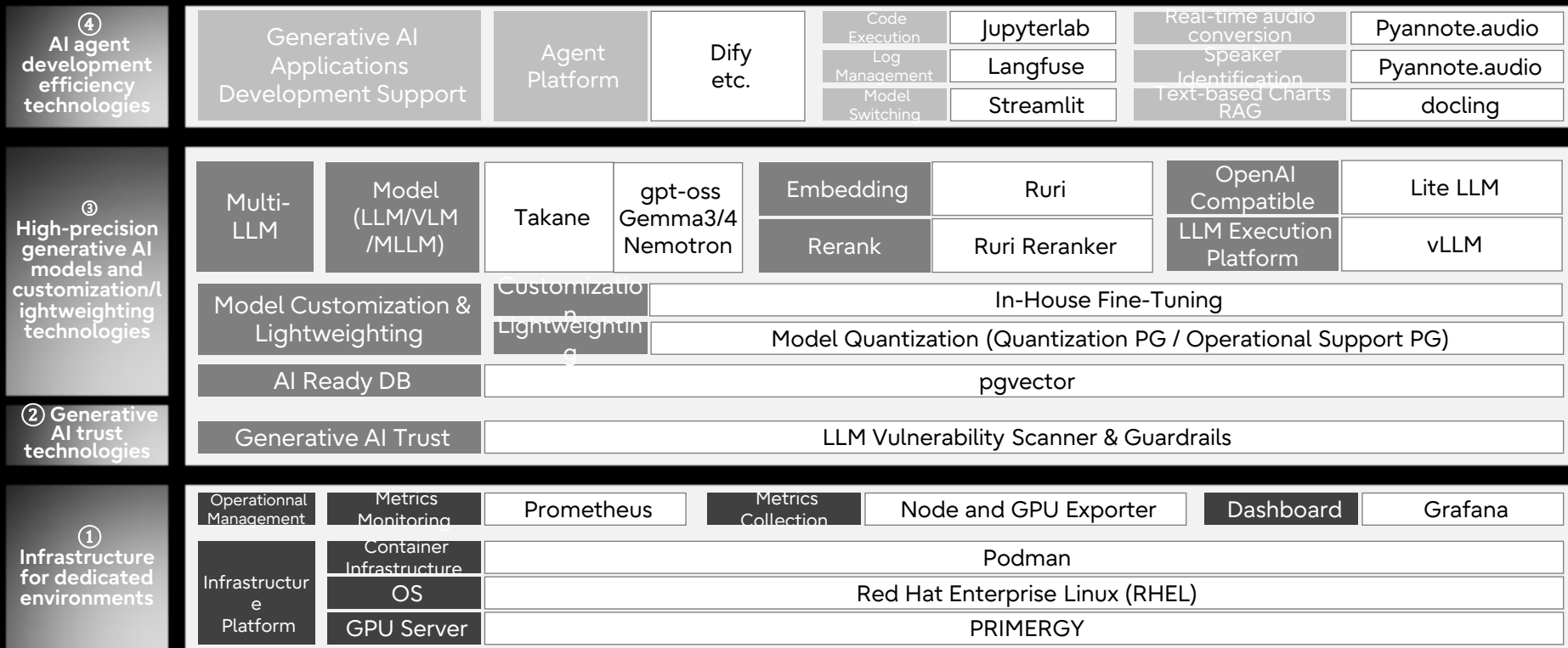
- Systemically managing administrative documents, operational manuals, and facility/operations information for iterative use across organizations.
- Sustainable operation based on closed and proprietary environments in the fields of disaster prevention, transportation, and energy, as well as national security and critical infrastructure domains.

# Fujitsu Kozuchi Enterprise AI Factory Details

# Function Stack (Configuration Overview)

uvance

- Providing all necessary functionalities for each layer (infrastructure, generative AI model, and application) in a single package.



# Private AI Platform on PRIMERGY[PAPP]

**uvance**



- **Pre-built Conversational Generative AI Platform**



A pre-built generative AI platform, finely balanced based on our AI engineers' practical knowledge, is delivered as a ready-to-use model.

- **Installation/Maintenance Services**



Reliable and smooth installation services by experienced engineers, along with hardware/OS maintenance services supported by our help desk, are provided.

- **Related Services**



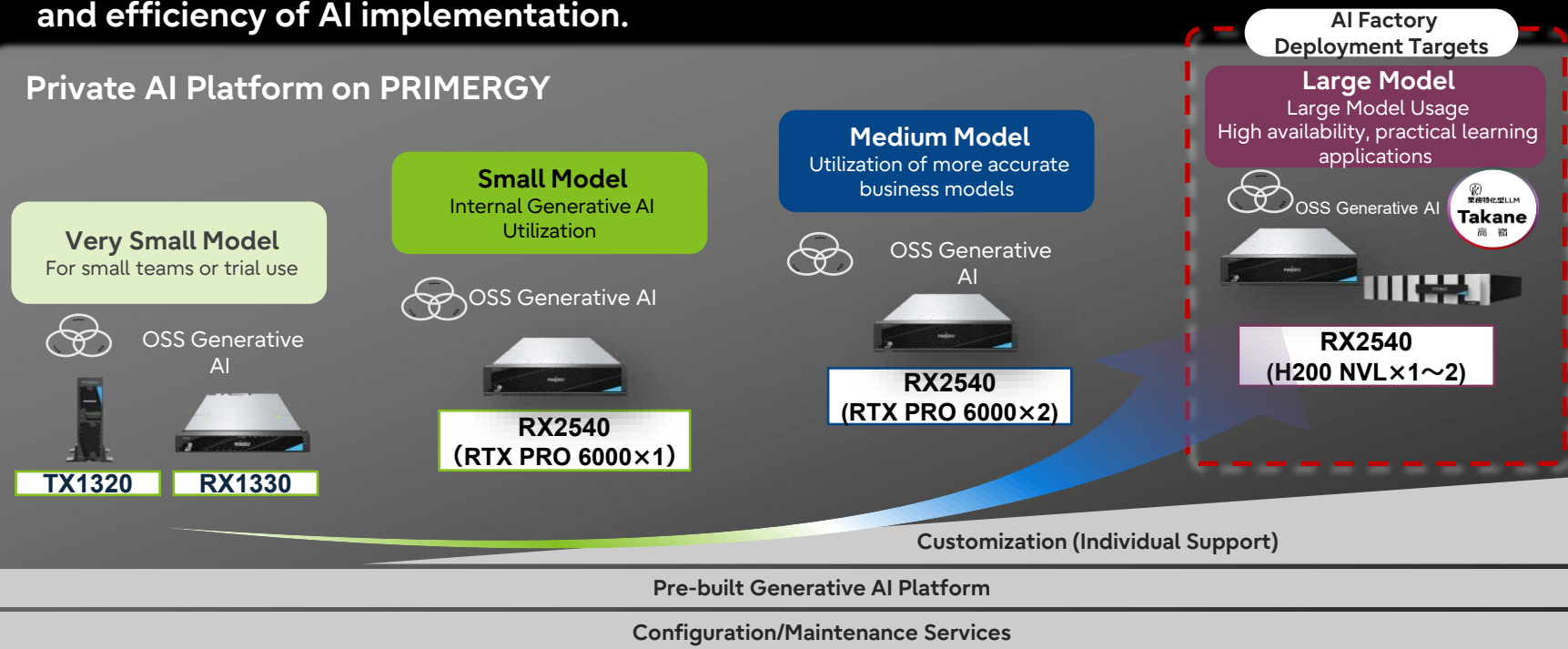
A variety of related services are offered to address diverse customer challenges, from pre-implementation assessment to post-implementation maintenance and operational support.

# Private AI Platform on PRIMERGY [PAPP]

uvance

- A diverse range of models, from Very Small to Large, is offered to suit various scales.
- Flexible configurations tailored to customer needs are possible, significantly enhancing the speed and efficiency of AI implementation.

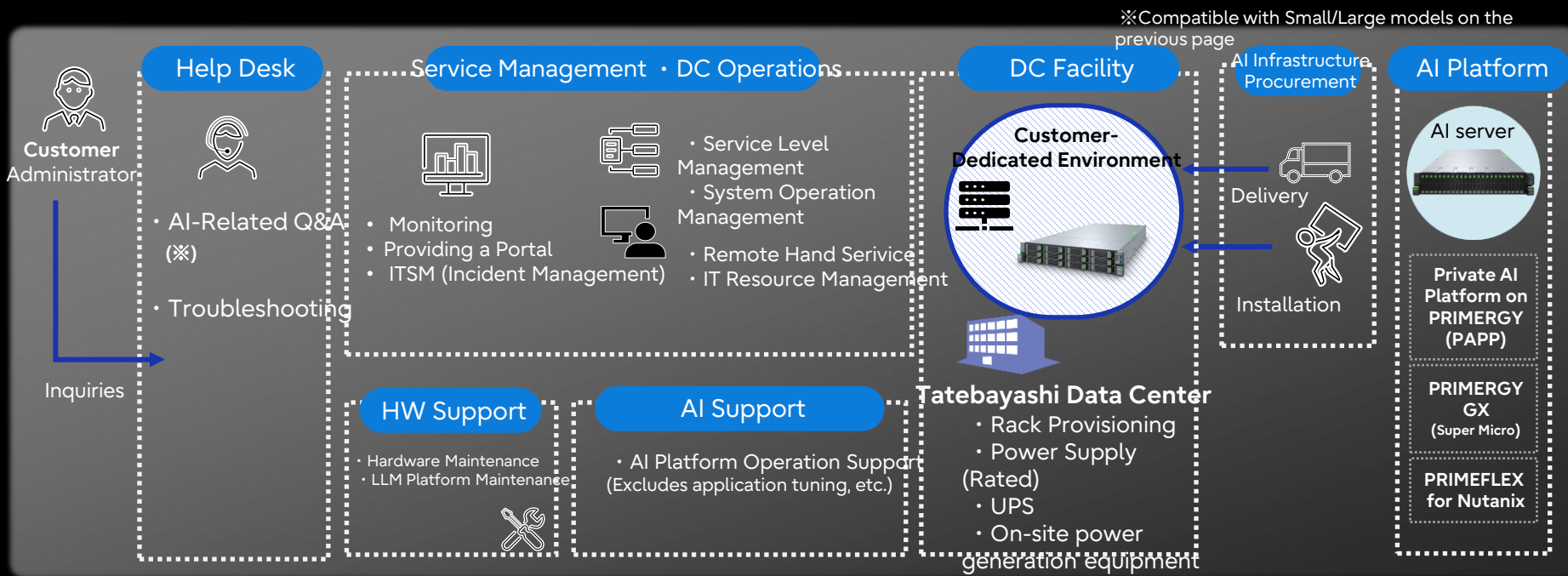
## Private AI Platform on PRIMERGY



# On-Premises Generative AI Solution [Managed]

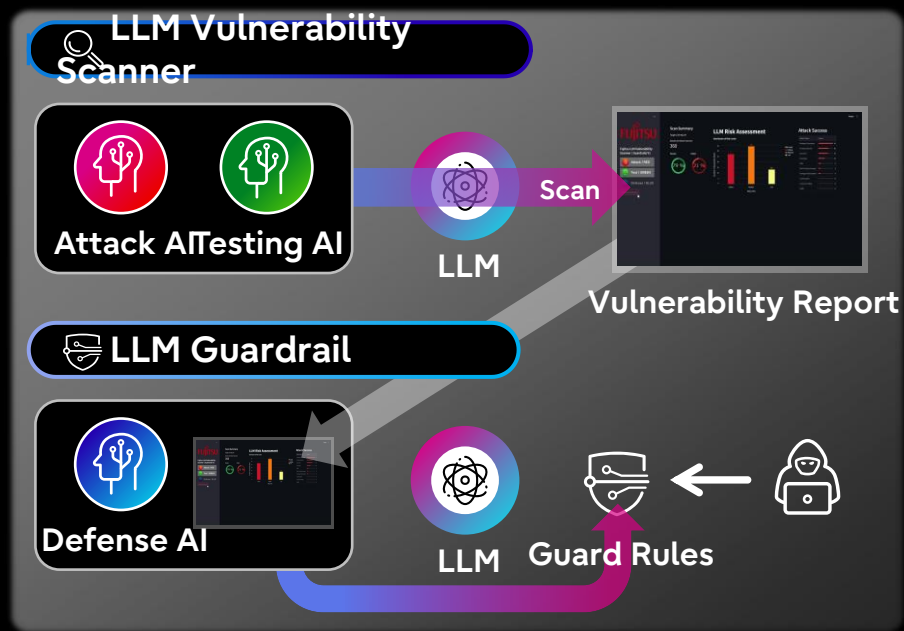
uvance

- Providing the necessary hardware (HW) and LLM environment for generative AI utilization, set up and delivered from Fujitsu's data centers.
- Configured as a dedicated environment, enabling use in a highly confidential setting without sharing with other companies.
- Offering full support with IT resource monitoring, covering troubleshooting upon malfunction and AI-related QA support.



# LLM Vulnerability Scanner & Guardrails

- LLM vulnerabilities are identified through comprehensive evaluation, and optimal defense technologies are automatically applied.



## Top-tier vulnerability coverage in the industry

- Supports over **9,000 types of vulnerabilities**, including Fujitsu's proprietary methods.
- Accurately analyzes **conversational attacks** that were previously difficult to detect.

## Requires no expert knowledge

- Explanation features enable **non-experts** to easily understand the impact.
- Enables enhanced security measures with **reduced learning costs**.

## Automatically generates defense rules

- Automatically generates** optimal defense rules based on reports.
- Enables out-of-band defense **without altering the core LLM**.

# LLM for Enterprise – "Takane"



- LLM optimized for Enterprise operations
- World-Class RAG Technology enabling high-precision use of in-house data
- Industry-leading Embed and Rerank models
- Highly customizable for key business tasks

Latest Large Language Model  
「Command R+」

- Japanese-specific additional training/Fine-tuning technology
- Composite AI technology
- Generative AI trust technology
- Provided in a private environment
- Knowledge accumulated through the building of mission-critical systems

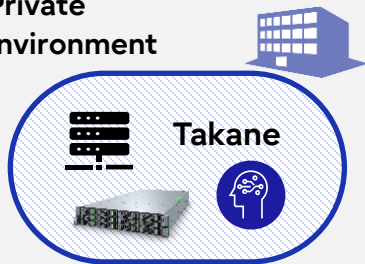
# Why Takane?

## ✓ Enterprise specific high-secure LLM

### Secure Data Utilization

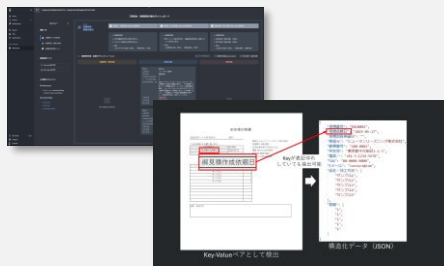
Confidential on-premises data that cannot be moved to the cloud can still be safely leveraged in a **closed environment**.

**Private Environment**



### AI That Knows Your Business

**Customize LLMs** to reflect your unique knowledge and deliver accurate, context-aware answers that strengthen competitive advantage.



### Reliable Implementation Support

**End-to-end assistance**—from strategy to operations—ensuring smooth adoption of generative AI with minimal customer burden.

Customer



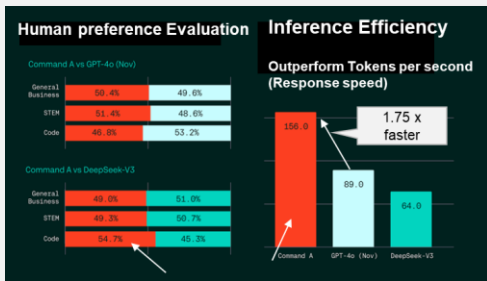
FUJITSU

# Why Takane?

- ✓ Delivers high Japanese performance with minimal infrastructure resources

## High-Speed Inference

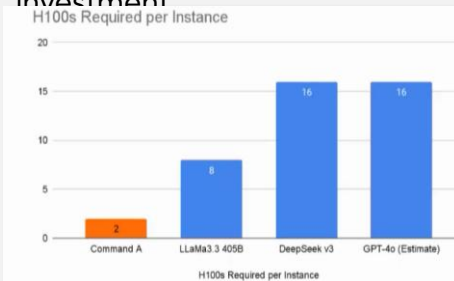
Delivers performance equivalent to general-purpose models while **achieving 1.75 times the output speed of GPT-4o**



Source: <https://cohere.com/blog/command-a>

## GPU-Efficient

When using NVIDIA H100, **operates with 1/8th the resources** compared to GPT-4o. Delivers maximum performance with minimal infrastructure investment



Source: <https://cohere.com/blog/command-a>

## Rich Lineup

Offering a diverse range of products from 1B to over 100B parameters

※2025.9 Company research

| parameter size | Cohere/Takane                           | NTT Tsuzumi                                  | NEC cotomi | IBM granite                      | Ricoh                | H2O                                                 |
|----------------|-----------------------------------------|----------------------------------------------|------------|----------------------------------|----------------------|-----------------------------------------------------|
| 1-30B          | Command R7B<br>Takane 2.0-7B            | Tsuzumi 0.6B<br>Tsuzumi 7B                   |            | Granite 3.3 2B<br>Granite 3.3 8B |                      | Denise30.5B<br>Denise30.5B 3.3B<br>Denise30.5B 3.3B |
| 30-100B        | Command R<br>Takane 2.0-12B             | Tsuzumi 12B<br>cotomi FastV2<br>cotomi ProV2 |            |                                  | 70B<br>Llama 3.3 70B | Denise34B                                           |
| 100B+          | Command R+<br>Command A<br>Command A MM | Takane 1.1<br>Takane 2.0-12B                 |            |                                  | gpt-oss-120B         |                                                     |

Differentiation

※1 As of November 2024, it is being evaluated as the medium-sized version of Tsuzumi.

※2 The parameter size has not been disclosed.

※3 In addition to the language model, Granite also provides vision and code models.

# Takane Lineup

- Select the appropriate model size based on your requirements

| Scale             | LLM Products       | Features                                                                                                                                         | Use Cases                                                                                                                                                 |
|-------------------|--------------------|--------------------------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------|
| Small             | Takane 7B          | <ul style="list-style-type: none"><li>• Small-scale Model</li><li>• Quick Response</li><li>• For small to medium-sized tasks</li></ul>           | <ul style="list-style-type: none"><li>• Enterprise Chatbot</li><li>• Content Generation</li><li>• Automatic Translation</li></ul>                         |
|                   | Command R7B        |                                                                                                                                                  |                                                                                                                                                           |
| Middle            | Takane 32B         | <ul style="list-style-type: none"><li>• Medium-scale model</li><li>• Contextual Understanding</li><li>• Complex Task Processing</li></ul>        | <ul style="list-style-type: none"><li>• Research Institutions</li><li>• Business Analysis</li><li>• Advanced Customer Support</li></ul>                   |
|                   | Command R          |                                                                                                                                                  |                                                                                                                                                           |
| Large             | Takane             | <ul style="list-style-type: none"><li>• Large-Scale Models</li><li>• Advanced Language Processing</li><li>• Specialized Domain Support</li></ul> | <ul style="list-style-type: none"><li>• Large-Scale Media</li><li>• Specialized Fields</li><li>• AI Assistant Development</li></ul>                       |
|                   | Command A          |                                                                                                                                                  |                                                                                                                                                           |
| Large Multi Modal | Takane 112B Vision | <ul style="list-style-type: none"><li>• Multi-modal Model</li><li>• Reading Charts and Images</li></ul>                                          | <ul style="list-style-type: none"><li>• Text Analysis Including Charts and Graphs</li><li>• Reading Drawings in Design and Manufacturing Fields</li></ul> |
|                   | Command A Vision   |                                                                                                                                                  |                                                                                                                                                           |
| Extension Model   | Embed 4            | • Maps text to a vector space, enabling similarity and semantic analysis                                                                         |                                                                                                                                                           |
|                   | Rerank 3.5         | • Reranks search results to quickly deliver highly relevant information                                                                          |                                                                                                                                                           |

## Diverse compatible models

### OpenAI gpt-oss 20b & 120b

OpenAI's high-performance  
Open-source LLM

A large language model that combines long input processing and complex contextual understanding, enabled by its massive parameters, with the flexibility for commercial use.

- gpt-oss-20b has 21 billion parameters, excelling in long-text processing and complex contextual understanding.
- gpt-oss-120b features 120 billion parameters, offering an accuracy comparable to GPT-4 while achieving both efficiency and precision.

### Google gemma3 12b & 27b gemma4 26b & 31b

A multimodal LLM developed by  
Google

A lightweight, high-performance open-source AI model that can process images in addition to text.

- Generate text using not only text input but also image input.
- Gemma 3 and 4 can process images, while Gemma 4 can process videos as frame-by-frame images.
- Supports over 140 languages, including Japanese.

### NVIDIA Nemotron3 nano

A new LLM model developed by  
NVIDIA

A high-precision LLM that supports complex decision-making through long-context understanding and advanced reasoning.

- Supports long texts of up to 1 million tokens while delivering high-speed inference.
- By combining the Mamba architecture, which maintains efficiency even with long texts, with the MoE architecture, which reduces computational requirements, the model achieves significantly higher efficiency than previous models, further improved by technologies such as FP8-based quantization, GPU optimization, and NVFP4)

# In-house Fine-Tuning & Model Quantization

uvance

## In-house Fine-Tuning

- Low response accuracy
- Requires resources for business-specific customization
- Lack of personnel with specialized expertise

Provides efficient  
tuning methods  
and execution  
environment

Enables customers to fine-tune the model on their own

If a training dataset is available  
Fine-tuning is possible



Fine-Tuning Function / Execution Environment

Enables the creation and improvement cycle of domain/business-specific models

## Model Quantization

Generative AI consumes vast amounts of electricity, so it needs to be made lightweight according to its intended use

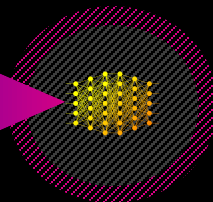
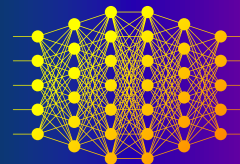
Achieved the world's  
highest accuracy with  
1-bit quantization

**Quantization:** A technique that compresses the weights assigned to connections between neurons  
(Reduces model size by representing weights with fewer bits)

Achieving lightweight and low power consumption  
while maintaining accuracy

Before quantization  
(32-bit or 16-bit)

After quantization  
(1 bit)



Maintains high precision using Fujitsu's proprietary quantization error propagation method

# Multi-AI Agent Framework

- Domain-Specific, High-Quality, Secure Workflow Automation for AI Agents

## Multi AI Agent Framework

Building



AI Agent

Connecting



Workflow

Orchestration



Control



Domain-Specific



Quality

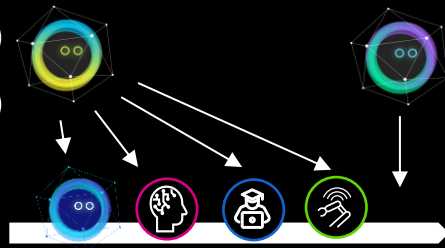


Security

Domain-Specific

Agentic Memory

Integration



Workflow

Monitoring

Constraint-Based Routing

Growth (Quality)

Acquires knowledge from humans, robots, and other agents.  
Execute workflows while growing within specific domains.

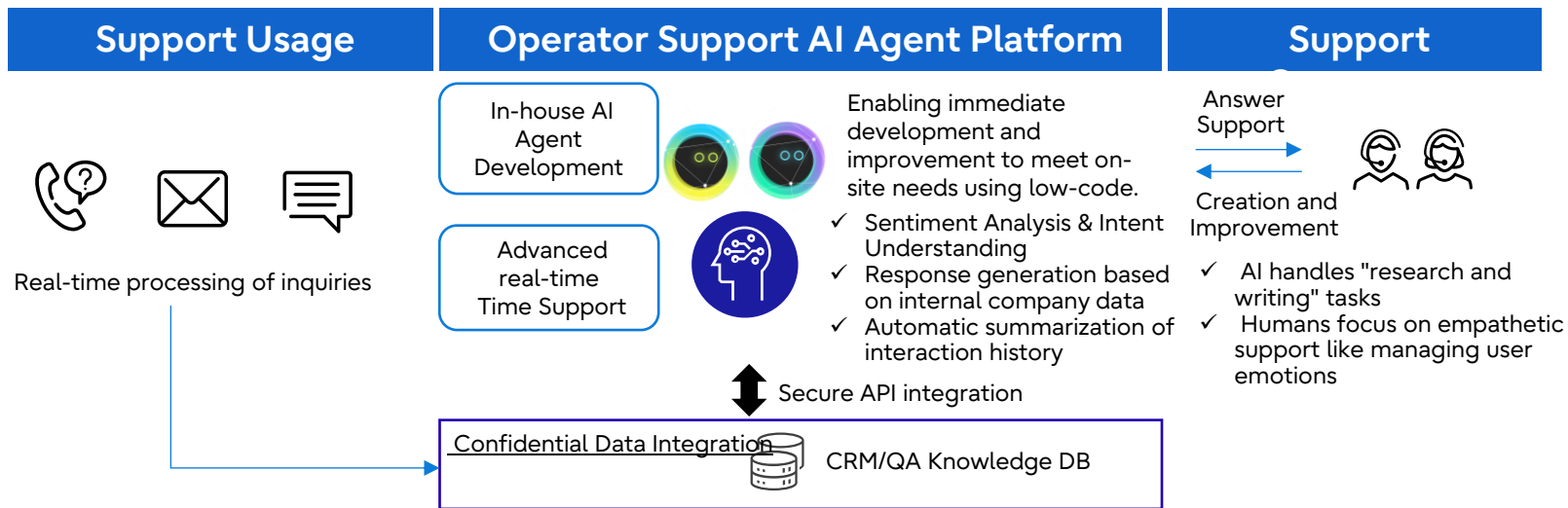
Properly monitors complex interactions throughout the entire workflow.  
Supporting workflow execution while considering diverse constraints.

# Use case

## Background / Purpose

- Building a next-generation contact center platform for secure, low-code, field-initiated AI agent development

## Utilization Usage Scenario



## Expectations Effect

- Improving QCD (Quality, Cost, Delivery): Enhancing operator response quality and reducing handling time through the automation of search and recording tasks
- Business Transformation: Shifting support operator responsibilities from considering and creating responses to developing/improving agents and providing empathetic support to inquirers.

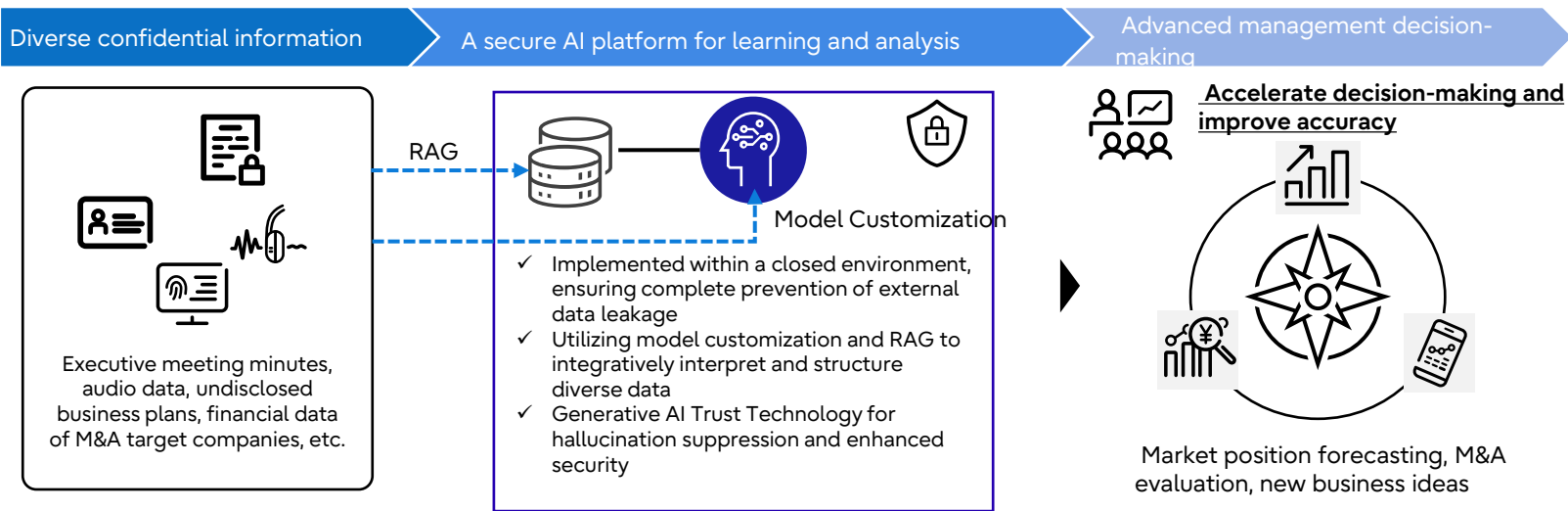
# Management Decision Support

uvance

## Background/ Purpose

- Extracting information from internal confidential documents and data sources to support executive decision-making

## Application Usage



## Expectations Effect

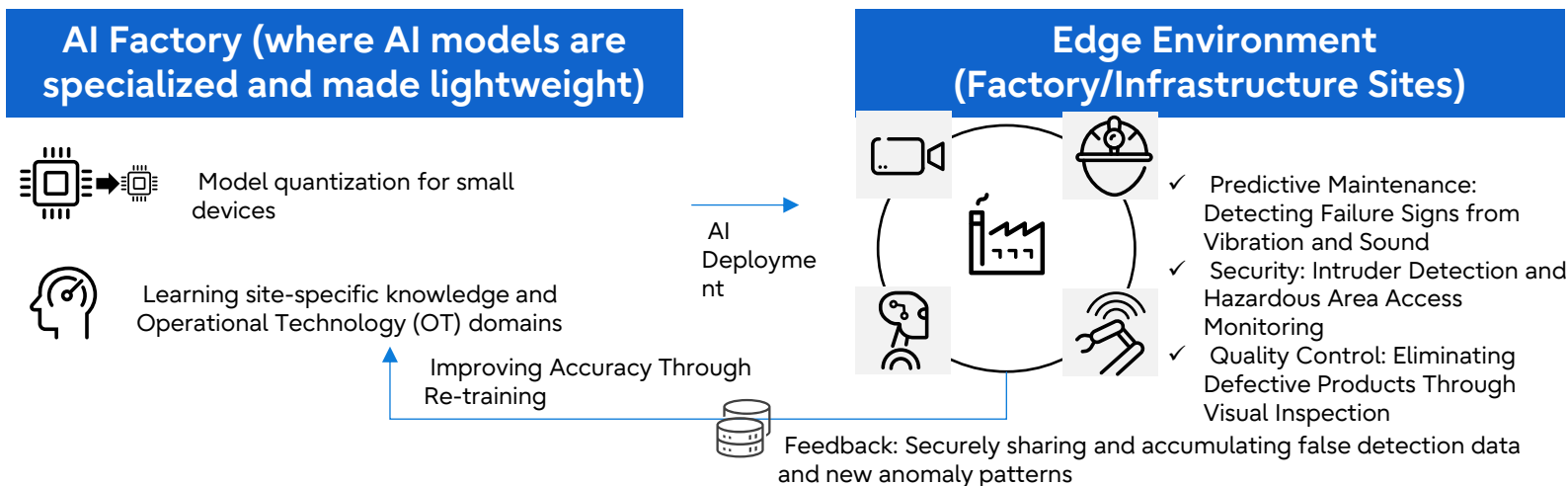
- Leveraging, analyzing, and inferring from data that can only be handled in a private environment, AI highly supports executive decision-making

# Edge AI evolving driven by manufacturing and infrastructure sites

## Background/ Purpose

- Deploying lightweight models via quantization technology and in-house fine-tuning to cultivate AI specifically tailored for on-site operations

## Application Image



## Expectations Effect

- Enhancing equipment utilization through predictive maintenance, automating visual inspections for consistent quality, and ensuring continuous monitoring of hazardous areas

# Pre-release trial

# Early Trial Announcement

Prior to the official launch, we will begin accepting registrations for an early trial starting February 2nd, which will offer access to selected features such as in-house fine-tuning and model quantization.

We plan to roll out services incrementally from February, with the official launch scheduled for July.

# Trial Specifications

|                                |                                                                                                                                                                                                                                                                                                                 |
|--------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Name                           | Fujitsu Kozuchi Enterprise AI Factory Early Trial                                                                                                                                                                                                                                                               |
| Trial Overview                 | A free trial environment is available, allowing users to experience a limited-feature version of Fujitsu Kozuchi Enterprise AI Factory, scheduled for official release in July 2026.                                                                                                                            |
| Target Audience                | Customers, Fujitsu Group Companies                                                                                                                                                                                                                                                                              |
| Trial Environment and Features | <ul style="list-style-type: none"><li>• Private AI Platform on PRIMERGY</li><li>• Fujitsu Generative AI for Cohere (Takane)</li><li>• In-house fine-tuning functionality</li><li>• Vulnerability Scanner/Guardrail Function</li><li>• LLM Quantization Function (*1)</li><li>• AI Agent Framework(*2)</li></ul> |
| Pricing                        | Free of charge for environment usage (service support, etc., is generally chargeable).                                                                                                                                                                                                                          |
| Support                        | Email-based Q&A                                                                                                                                                                                                                                                                                                 |
| Usage Period                   | The usage period will be determined during the preliminary consultation. (Estimated: 2 to 4 weeks)                                                                                                                                                                                                              |
| Application Start              | February 2, 2026                                                                                                                                                                                                                                                                                                |
| Notes                          | After conducting preliminary interviews and assessments, we will select the appropriate trial environment and features and determine their availability. As the number of available trial environments is limited, we may not be able to accommodate all scheduling requests.                                   |

(\*1) The service provider will perform the quantization of Takane and provide it. The functionality allowing users to quantize LLMs themselves will be added in phases.  
(\*2) Only, standard on the Private AI Platform on PRIMERGY, is available for use. AI Agent functionality will be added in phases.

# Steps to Start Your Trial

**uvance**

## STEP 1 Inquiry

- ✓ For inquiries, please contact the trial desk through your assigned sales representative.

## STEP 2 Assessment

- ✓ We will conduct interviews regarding your planned evaluation content, features to be used, usage schedule, and prospects for future full-scale utilization.
- ✓ Based on the interview results, we will determine the feasibility of the trial and adjust the schedule.

## STEP 3 Application

- ✓ Formal application will proceed based on the assessment outcomes.

## STEP 4 Environment Provision

- ✓ We will set up and provide the environment.

