

Self Evolving Specialized LLMs on HPC: Efficiency via 1 bit Quantization and AI Computing Broker, Toward NPUs

Koichi Shirahata

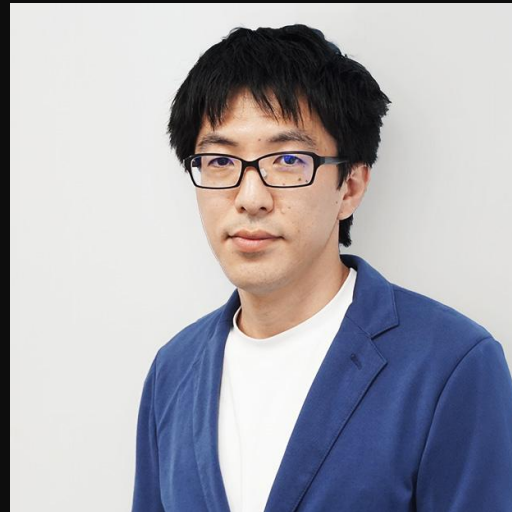
Fujitsu Limited

June 24, 2026

Koichi SHIRAHATA

Senior Project Director - Artificial Intelligence
Laboratory, Fujitsu Research, Fujitsu Limited

- March 2015: Received Ph.D. at School of Information Science and Engineering, Tokyo Institute of Technology
- April 2015: Joined FUJITSU LABORATORIES LTD.
- November 2020, 2021: Achieved the world's highest performance in MLPerf™ HPC, a machine learning performance benchmark, using the Fugaku supercomputer and ABCI
- May 2024: Development of a distributed parallel training method for large-scale language models in the policy response framework of the supercomputer "Fugaku"



Fujitsu AI Strategy



Providing innovative **sovereign** AI platforms for enterprise customers



Manufacturing



Defense



Healthcare



Government



Finance

Security

Computing

FUJITSU-MONAKA
Quantum

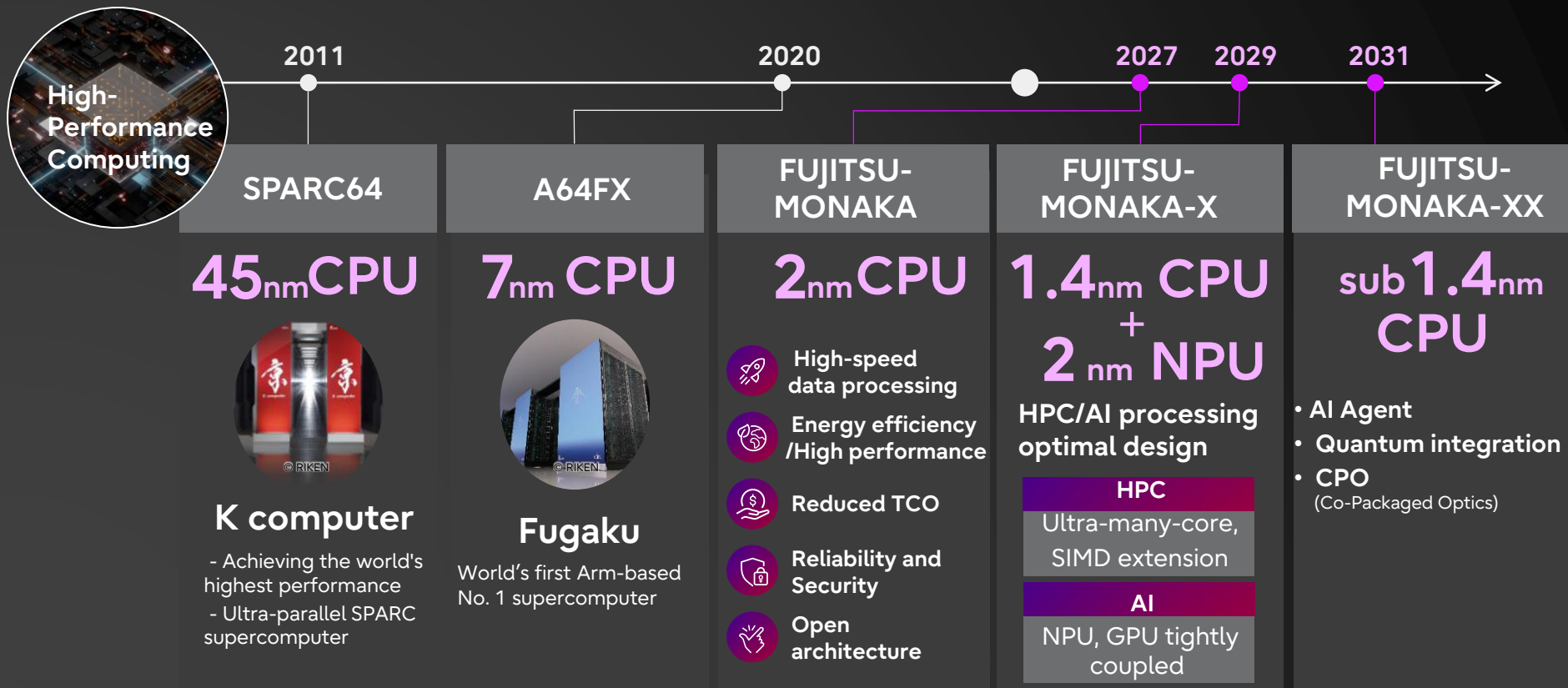
Network

Photonics
Wireless

AI software

Fujitsu Kozuchi
Kozuchi Physical AI
Takane

Computing for AI – FUJITSU-MONAKA(CPU/NPU)



Fujitsu Kozuchi AI Platform for Sovereign AI

FUJITSU

Transforming business processes and corporate culture to achieve high productivity and employee satisfaction

Domain-Specific

Generating AI technology reflecting each company's unique strengths

Flexibility

Optimizing self-evolving learning, and large-scale agent collaboration across organizations

Security

Addressing evolving LLM vulnerabilities and ensuring confidentiality and safety in cross-organizational agent collaboration



Kozuchi AI Security Platform

Addressing new cyber threats



Predictive Resilience

Anticipating the future and proactively addressing emerging risks

AI Scanner/ Guardrails

Multi-agent Security

Secure inter-agent gateway

Creating Safety and Security in the Information Society



Self-Evolving Multi-AI Agent Technology

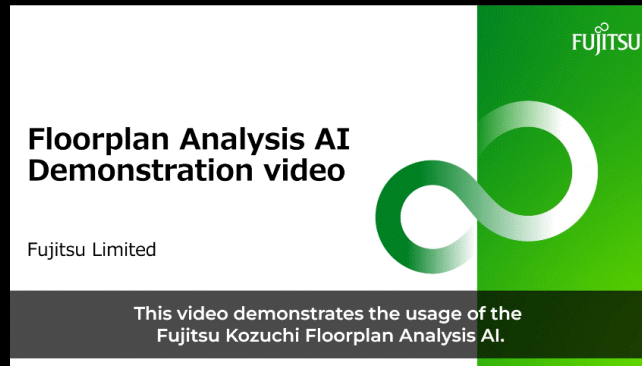
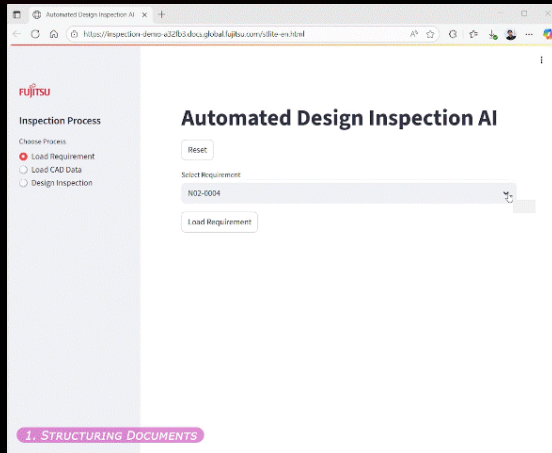
Challenges in fields such as manufacturing and construction are difficult to solve with general-purpose AI: Fujitsu has developed "Takane," a business-specific large language model



Manufacturing



CAD Automatic Inspection



Construction



Floor Plan Analysis

Problem

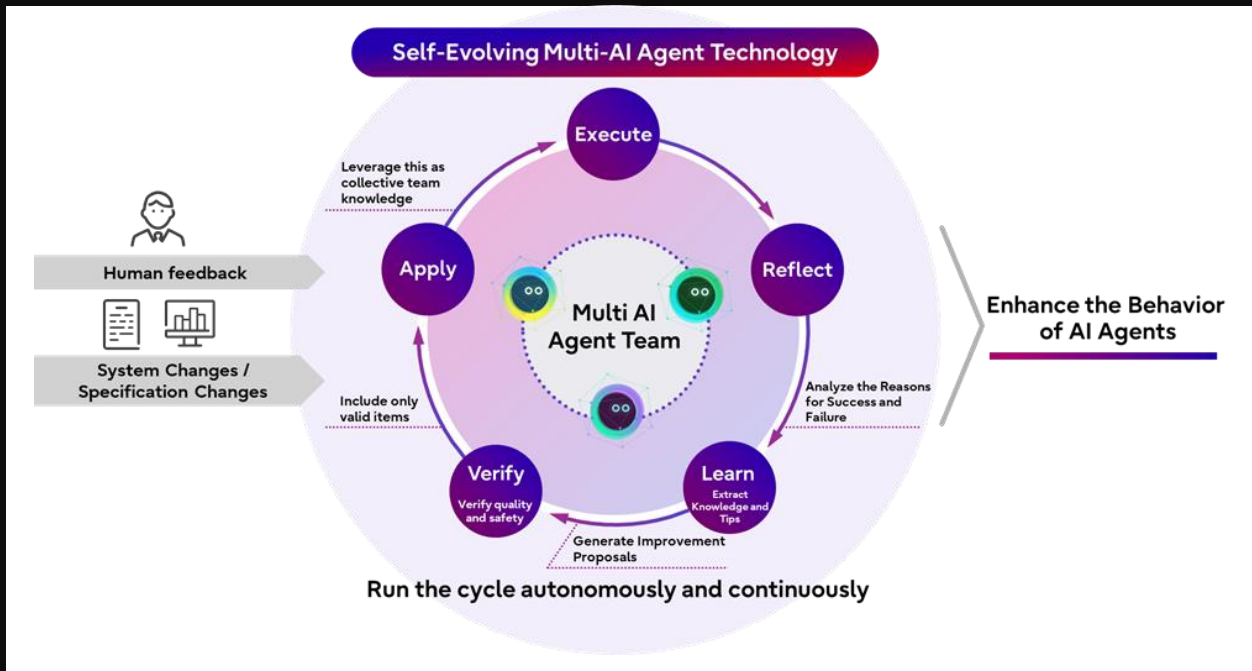
Because specialized design requires expert knowledge, there are limits to the precision that can be achieved by human designers.

Model training, evaluation, and improvement take too much time, making it impossible to roll out the initiative across the organization

We have developed technology that autonomously generates and evolves specialized AI through a cycle of operational analysis and learning

Self-Evolving Multi-AI Agent Technology

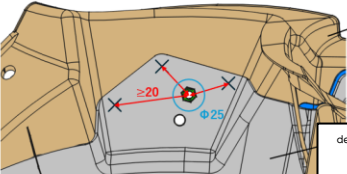
- Enabling automatic enhancement of the business-specific LLM "Takane" and safe self-evolution of AI agents
- Enables multiple AI agents to perform tasks as a team, continuously and safely learning from daily execution results, human feedback, policy revisions, and specification changes.



Inspection Tasks in Manufacturing Case Studies

Check the conformity of CAD data using a test program generated from the requirements specification

Requirements Document

Requirement ID	L02-0001	Requirement Name	Constraints on the installation locations of hexagonal rivets for the rear door interior panel reinforcement	Remarks		Scope of Disclosure	
Part Name	Rear Door	Category	Assembly Ease	Rank	Required	Target	All
Ensure a flat surface with a diameter of 25 mm at the installation points for the hexagonal rivet nuts on the rear door inner panel reinforcement. In addition, ensure that the distance from the hex rivet nut to the spot weld is at least 20 mm. $\geq 20 \text{ mm}$ For the installation location of the hexagonal rivet nut on the inner panel reinforcement of the rear door, ensure a flat surface with a diameter of $\phi 25$. Additionally, maintain a distance of 20 mm or more from the hexagonal rivet nut to the spot welding point. Design example  Inner panel of rear door							

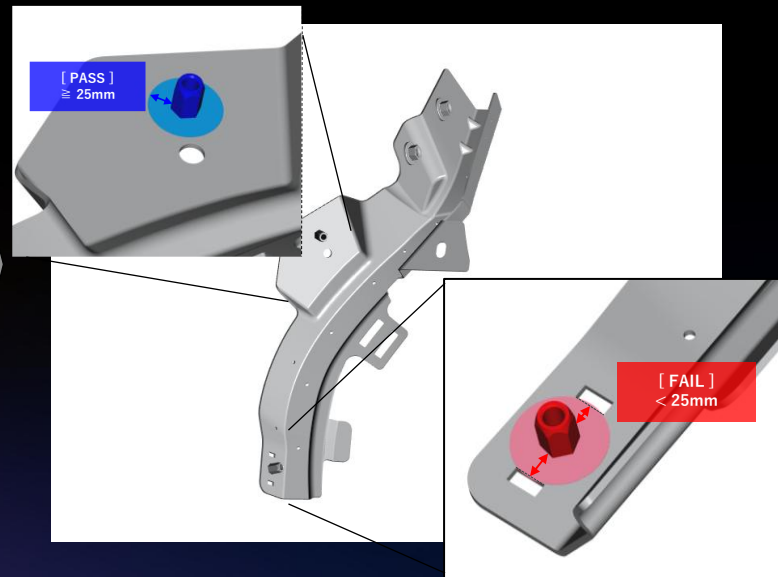
Inspection Program

```
def is_flat(surface, r_list):
    flat_results = []
    radius = 12.5 # 半径12.5mm

    for rivet_nut_info in rivet_nut_info_list:
        rivet_nut_center = np.array(rivet_nut_info["center_on_panel"])

        # 半径範囲内のメッシュのポイントを取得
        in_radius_indices = np.where(np.linalg.norm(surface_mesh.vertices - rivet_nut_center, axis=1) <= radius)[0]
        # 対応する頂点を抽出
        in_radius_vertices = surface_mesh.vertices[in_radius_indices]
```

CAD data to be inspected



TAKANE SELF-ADAPTING MULTI-AGENT SYSTEM

Self-Evolving Agents That Automatically Build Domain-Specific LLMs

Each agent understands the nature of the task and autonomously
adapts its strategy, methods, and evaluation criteria while cooperating
dynamically

Self-Evolving

Dynamic Adaptation

Fully Automated

Proven in 4 Domains

I



Title

Prompt

Pipeline

Inference

Expansion

Vision

Fujitsu

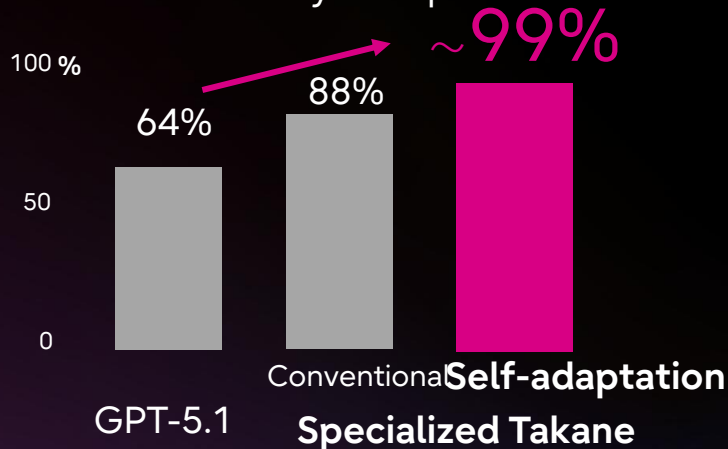
TAKANE - CONFIDENTIAL

Development Technologies and Customer Value

Self-evolving multi-
AI agent
technology

Dynamic multi-agent technology that continuously updates the weights of multiple agents through reinforcement learning, enabling the entire team to specialize and stabilize through interaction

Accuracy Comparison



Development time for specialized AI



Customer
Values

Automatically acquire "task-specific AI" with higher accuracy than human-created AI

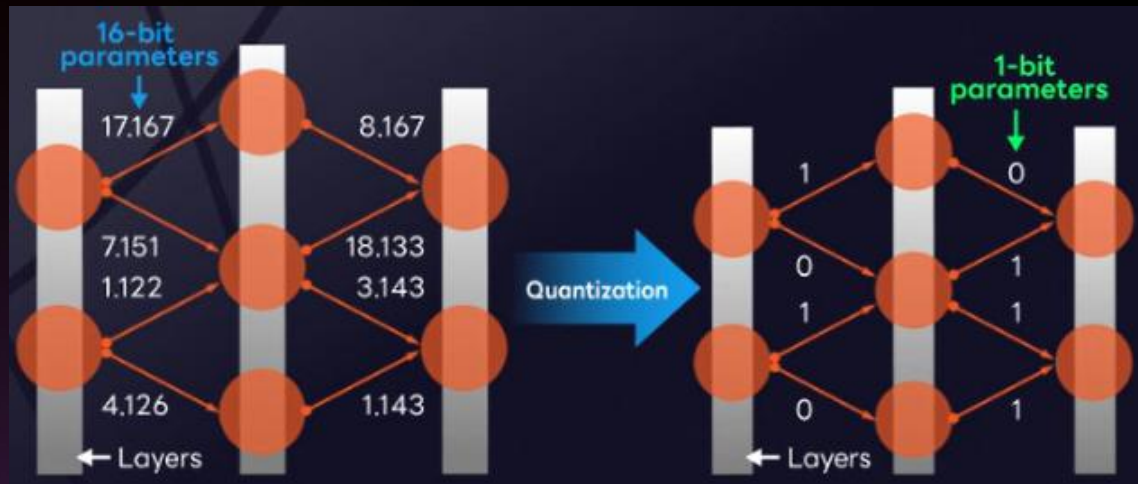
Not limited to a single task, but scalable to multiple tasks and locations

We can deliver a highly accurate, customized Takane tailored to customer needs in as little as one day to one week

1bit Quantization Technology

What is Quantization?

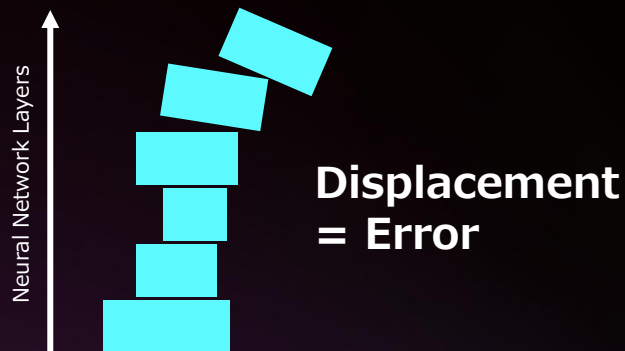
A technology that compresses floating-point operations within the AI brain "neural network" by replacing them with simpler forms, thereby reducing the memory and computational load of AI models.



Challenge: since the accuracy decreases as the number of bits is reduced at 1 bit, no practical application exists

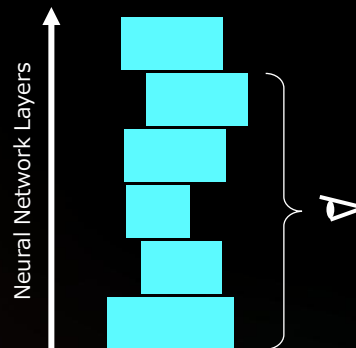
Our Developed technology

Conventional technology



Displacement occurs in each layer, and as the number of layers increases, the displacement grows larger and collapses (precision breakdown).

QEP (Quantization Error Propagation)



Error propagation method

A method for adjusting while checking for misalignment to ensure the entire structure stacks straight

Calibration Data Utilization

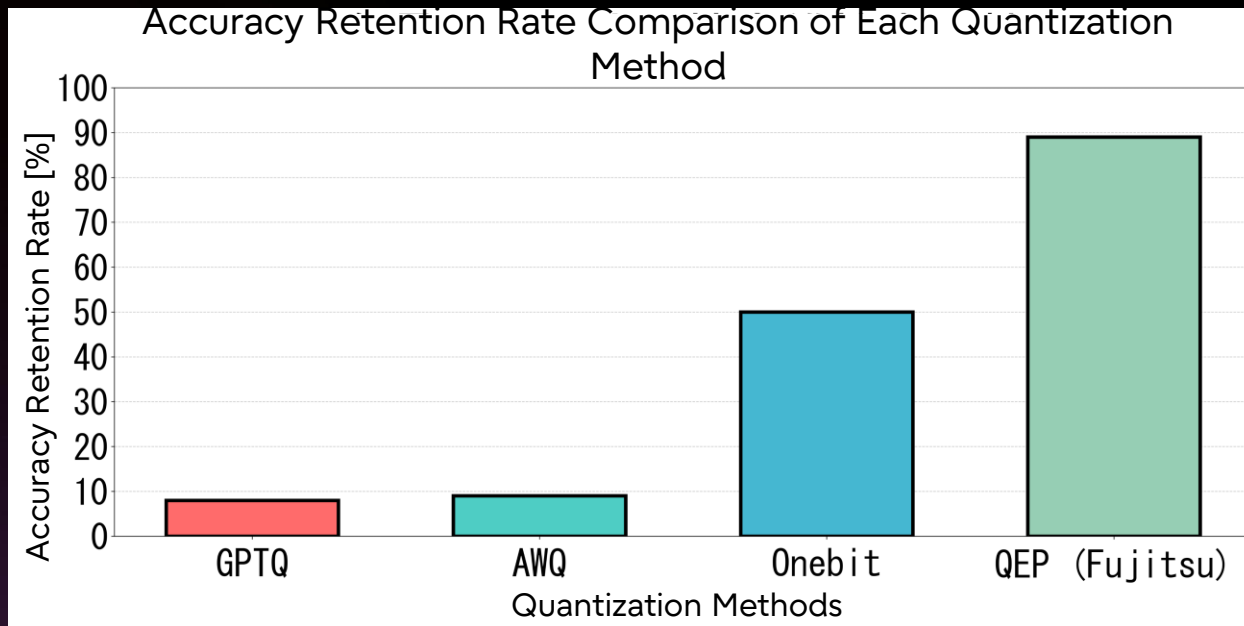
Use sample data to determine the stable position in advance.

QQA (Quasi-Quantum Annealing)

World-leading large-scale optimization algorithms developed through Digital Annealer

Achieved the world's highest precision retention rate of 89%

Accuracy Comparison with Conventional Technology



Achieved the world's highest accuracy retention rate of 89% with 1-bit quantization. Memory usage reduced by up to 94%, enabling code generation via 32B LLM to run at the edge

We have released the core functionality as open-source software: **FUJITSU**
OneCompression



<https://github.com/FujitsuResearch/OneCompression>

AI Computing Broker

Optimization for AI - Computing Workload Broker



A platform combining HPC and quantum computing technologies

- Configures a quantum/HPC hybrid framework tailored to the operational characteristics of each computational type
- The framework predicts CPU/GPU/quantum behavior and optimally allocates resources based on each characteristics

Application



AI

Simulation

Data Analytics

Platform

Computing Workload Broker (CWB)

Low power consumption

High computational performance

Reliable platform

Middleware

OS

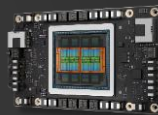
Hardware



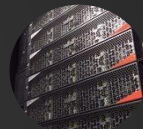
CPU
(MONAKA/-X)



GPU
NVIDIA, AMD



Quantum-inspired technology
Quantum simulator/Digital Annealer



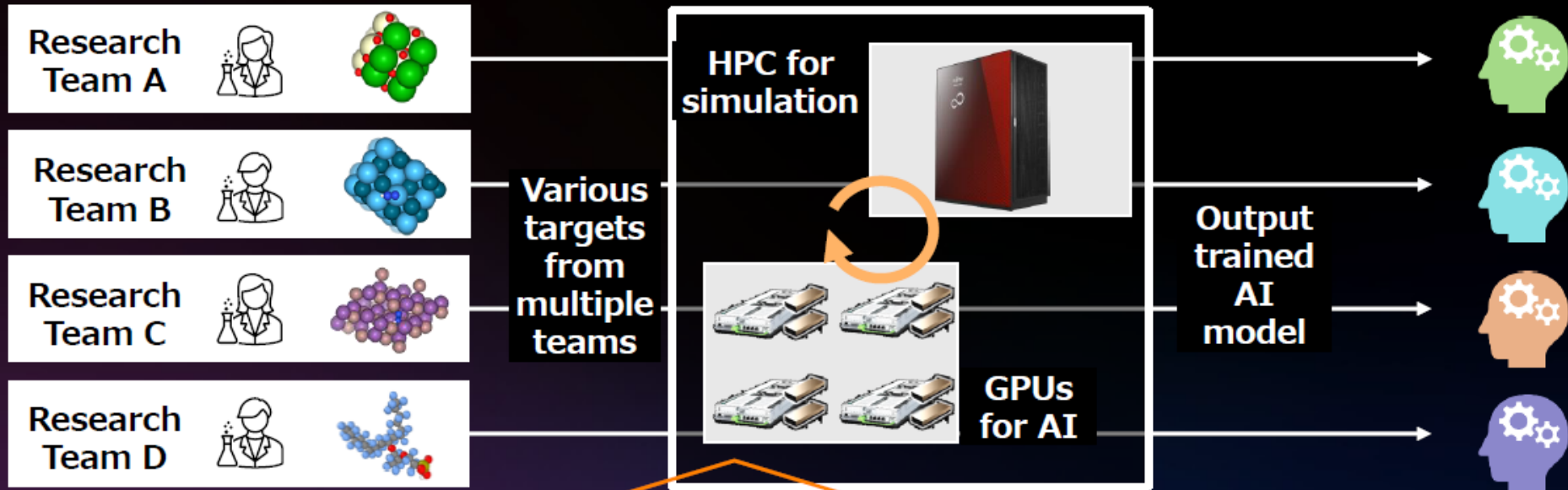
Quantum computer
Superconducting/Diamond spin



©RQC

Practical Use Case of NNP Generator

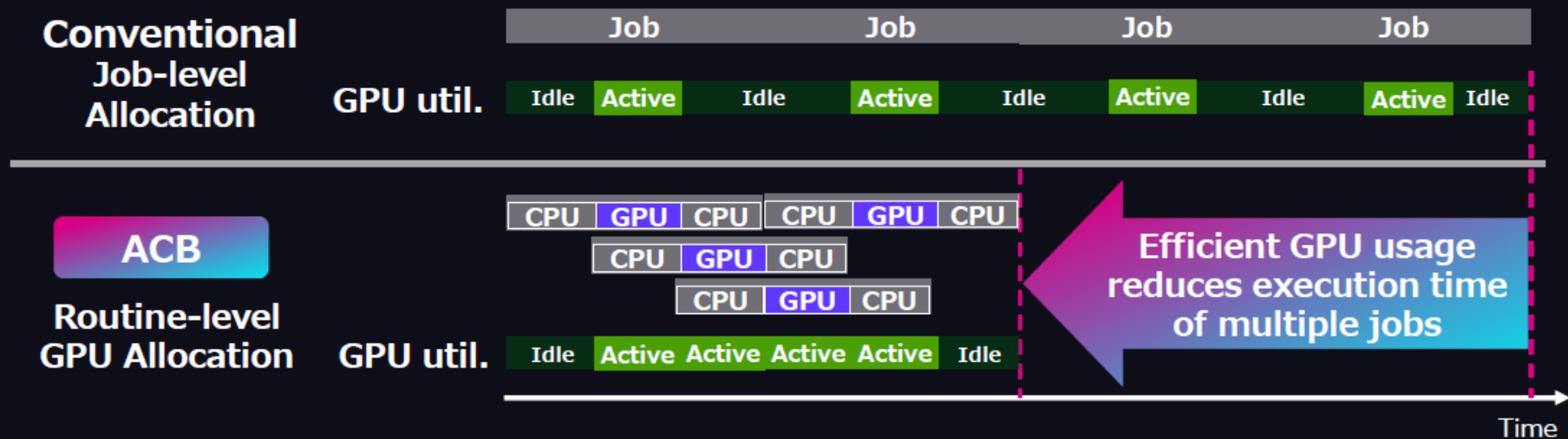
Multiple AI training for each target are running simultaneously on limited computing resource



Challenge: Heavy AI-training jobs monopolize GPU, delaying short jobs significantly

Best-in-class GPU utilization efficiency

- “Routine-level” allocation that detects actual GPU parts of jobs and dynamically allocates GPU accordingly



Enabling full GPU memory for each job

- Allocate GPU to only one job at a time (Temporal-sharing)
- Data of other jobs on GPU is automatically swapped to CPU

Conventional: Spatial-sharing

Memory is divided among jobs
Limited to small AI models



Memory

Job A

Job B

Job C

Computing Unit

Sharing among Job A/B/C

ACB

Time-sharing

Memory is occupied by each job
Large AI models can run



Memory

Job A

Comp. Unit

Job A

Memory

Job B

Comp. Unit

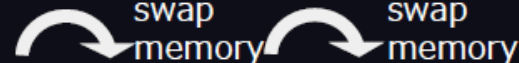
Job B

Memory

Job C

Comp. Unit

Job C



AI computing broker (ACB)

A middleware to share GPUs among AI apps.

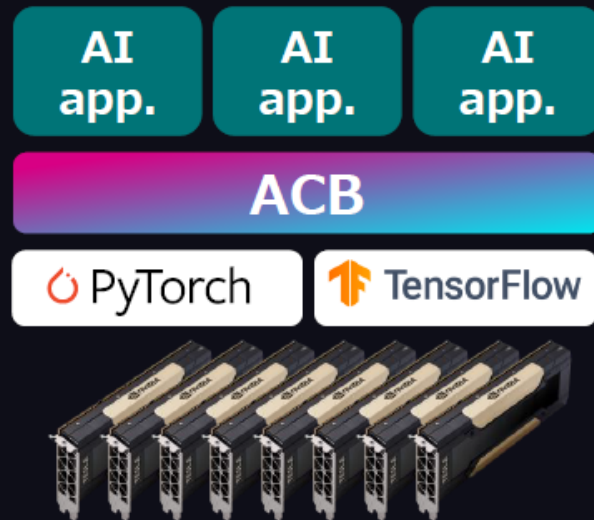
Key Features

- Best-in-class GPU utilization efficiency
- Optimized memory management across various AI applications

"Doubled model training throughput!"

"Deploying multiple jobs beyond physical GPU memory capacity!"

Success stories from ACB users
Scan here for more detail!



Key Takeaways

- AI is evolving from static models to self-evolving systems
- Fujitsu's self-evolving multi-AI agent technology automates the full lifecycle: data generation, training, and optimization
- Deployable efficiency technologies (1-bit quantization, distillation, AI Computing Broker (ACB)) enable AI from HPC to edge environments
- Looking ahead, these techniques will be integrated with emerging NPUs through hardware–software co-design
- Extending specialized AI into resource-constrained and sovereign environments

Thank you