



FOR IMMEDIATE RELEASE

Launch of the Liquid-Cooled GPU Container "DC-PKG": An AI Infrastructure Platform for Using Confidential Data in a Dedicated Environment

Design, build, monitoring, and operational support delivered as a single package

Sep 08, 2026 – Fixstars Corporation (TSE Prime: 3687, US Headquarters: Irvine, CA), a leading company in performance engineering technology, NTTPC Communications, Inc. (NTTPC), and Getworks Co.,Ltd. (Getworks) today announced the launch of the "Liquid-Cooled GPU Container DC-PKG," a single package that combines a container-type data center, NVIDIA accelerated computing servers, and integrated monitoring and operational support services.

As today's data centers adopt next-generation accelerated computing designed to scale with large AI models, cooling has become a major challenge. To address it, the three companies have conducted joint verification of liquid-cooled GPU server environments. By packaging those results together with each company's technology and expertise, the offering addresses the challenges of deploying and operating AI infrastructure.

► Press release on the joint verification (December 17, 2025) "Successful PoC on improving the operational efficiency of liquid-cooled GPU servers, advancing toward commercial use in Japan" (<https://www.nttpc.co.jp/press/2025/12/202512171500.html>)

1. Background

As AI adoption moves into full-scale operation, managing the underlying infrastructure has become a strategic priority

Corporate AI initiatives are moving from PoC to production, and agentic AI, generative AI, and AI inference are rapidly being embedded into real business processes. At the same time, while accelerated computing solutions continue to improve performance per watt, rising power consumption has made the shortage of data centers equipped with adequate cooling facilities a serious issue. Demand for deploying AI infrastructure on a company's own premises is also growing, driven by the need to protect confidential data, network constraints, the communication costs of large data volumes, and the operational burden of remote facilities. In building and operating such AI infrastructure, companies face the following challenges:

- Lack of design and construction expertise for the latest liquid-cooled accelerated computing servers
- Shortage of specialized personnel able to fully exploit accelerated computing server performance
- Growing burden of operations and maintenance after deployment

As a result, a growing number of companies find themselves wanting to use AI but unable to secure the infrastructure or the operational organization to do so. In some cases, this makes the path from evaluating AI to running it in production a long one, hindering swift progress toward stronger competitiveness and new business transformation.

2. The value of the Liquid-Cooled GPU Container DC-PKG

1) One-stop deployment of on-premises AI infrastructure on the customer's own site

Using the customer's land and power facilities, the package delivers a container-type data center, liquid-cooled GPU servers, and monitoring and operational support together as a single offering. This reduces the burden of vendor selection and coordination that arises at each phase of designing, building, and operating AI infrastructure. By bringing together the three companies' knowledge of container-type data centers, accelerated computing platforms, and AI software optimization, it shortens the time to deployment and enables a smooth start-up.

2) Stable and efficient use of high-performance accelerated computing servers

Integrated monitoring provides a single view of everything from container facilities to accelerated computing servers, enabling early detection of and response to interruptions. This limits the opportunity loss caused by server downtime and allows costly GPU server resources to be used efficiently. It also makes effective use of capital investment while holding down communication costs, optimizing operating costs over the medium to long term. Together, these benefits improve the return on investment in high-performance accelerated computing servers.



[Figure: Components of the package]

3. Anticipated use cases

1) Confidential Data Processing

Meets strict security requirements for data and networks, enabling AI analysis and generative AI use without moving confidential data outside the organization — achieving both security and more advanced operations.

2) Manufacturing and physical AI

Enables production use of visual inspection and anomaly detection AI based on design data and video from production lines, improving quality and streamlining inspection work. A company's own land and power facilities can also be put to use.

4. Package overview

Built around liquid-cooled NVIDIA accelerated computing servers and a container-type data center, the package integrates the components that make up an AI infrastructure platform.

Container-type data center	40-foot container; CDU chiller (liquid cooling unit for GPUs); ceiling-mounted and InRow air conditioning (temperature and humidity

	control inside the container); surveillance cameras (security monitoring); other facility equipment
Hardware / middleware	NVIDIA accelerated computing servers (configurations include HPE ProLiant Compute XD685 with NVIDIA HGX™ B300); firewall/UTM (security appliances); switches (network connectivity)
Monitoring and operational support	Performance visualization, integrated monitoring, operational support

* Items not listed above are outside the scope of this package.

* Construction work required to install the container, such as civil works and power supply, is not included in the package. These services are available separately upon request.

5. Pricing

The following are reference prices. Detailed pricing is quoted individually.

Configuration	Initial fee (excl. tax)	Monthly fee (excl. tax)
Small configuration (1 GPU server)	From ¥300 million	From ¥750,000

* Estimated prices as of August 31, 2026.

- Lead time: from 8 months * Lead times assume a standard configuration and are indicative only. They may be extended depending on the customer's site conditions, installation environment, equipment procurement, and other factors.
- Customization: Components can be reconfigured flexibly to match customer requirements. Please contact any of the three companies for details.

6. Availability

September 8, 2026 (Tuesday)

7. Inquiries

Please direct inquiries to the following URL.

https://dm.nttpc.co.jp/form/inq_ai_solution.html

For package details, please see the service website.

<https://www.nttpc.co.jp/ai-sol/liquid-cooling/>

8. Future plans

Going forward, the three companies will work to support next-generation accelerated computing servers (such as the NVIDIA Vera Rubin platform) and to expand into distributed AI infrastructure and edge AI, further broadening the infrastructure services that underpin corporate AI adoption. Through these efforts, they will respond flexibly to advances in AI technology and to wider adoption, supporting more sophisticated use of AI and sustained business growth for their customers.

Company profiles

About NTTPC Communications, Inc.

As the ICT services expert for mid-sized and small businesses in Japan within the NTT Docomo Business Group, NTTPC provides a range of services that are simple and priced for easy adoption.

The company is actively engaged in AI-related businesses, including the construction of GPU infrastructure, as well as in its network and cloud/data center businesses.

Its book *GPU Clusters × Generative AI: A Practical Guide to Next-Generation Infrastructure and Visualization in 13 Points* offers an accessible explanation of technologies such as building generative AI infrastructure.

► <https://www.nttpc.co.jp/press/2025/06/202506131500.html>

About Getworks co.,ltd.

Getworks announced its first container-type data center in 2013 and has continued to develop the product through a wide variety of demonstration projects. Beyond simply building data centers in containers, the company works with local governments on energy efficiency and renewable energy, making use of various renewable sources such as snow, water, and outside air.

Continuing to build configurations suited to the environment of each site (climate, power supply, and so on) and to customer needs, Getworks had a track record of 300 units built as of the end of January 2026 (270 20-foot and 30 40-foot units), including deliveries to customers with demanding requirements such as major corporations, power utilities, and

hospitals. The company can handle everything from site selection to civil, electrical, and telecommunications work and the associated applications. It also supports various subsidies and grants, which a growing number of customers have been using in recent years.

Driven by growing demand for AI and high-speed computing, Getworks has installed and operated more than 3,000 servers and over 10,000 GPUs. Fully in-house design and domestic production deliver short lead times and lower costs, with delivery and start of operation possible in as little as 10 days from order.

About Fixstars Corporation

Fixstars is a technology company dedicated to accelerating AI inference and training through advanced software optimization solutions. It supports innovation in healthcare, manufacturing, finance, mobility, and other industries. For more information, visit: <https://www.fixstars.com/>

* NVIDIA and NVIDIA HGX are trademarks or registered trademarks of NVIDIA Corporation in the United States and other countries.

* HPE and other trademarks used by HPE are registered trademarks of Hewlett Packard Enterprise Company and/or its affiliates.

###

Media Contact

Public Relations, Fixstars Corporation

Email: press@fixstars.com

Tel: +81-3-6420-0751