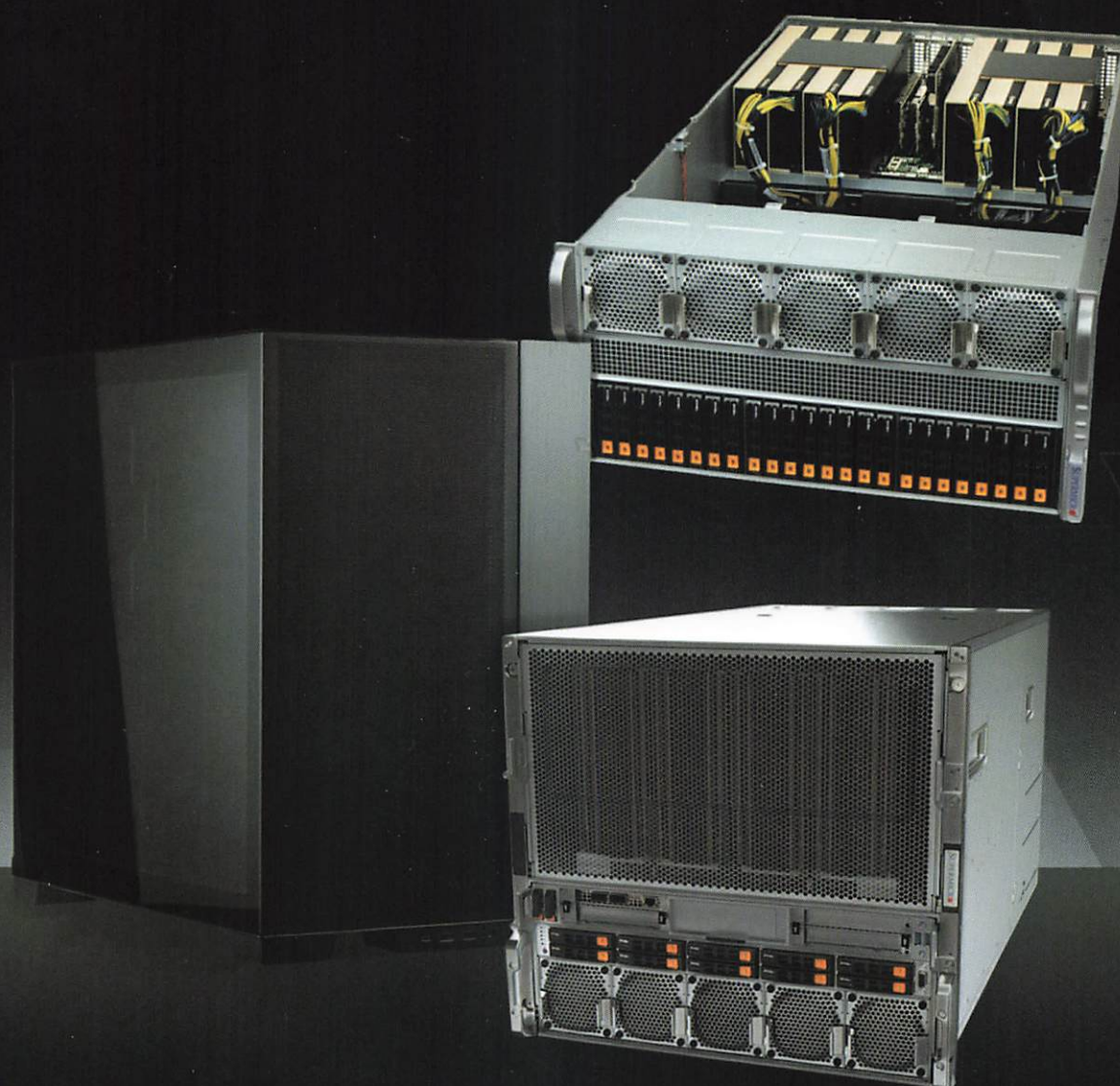


GENERAL CATALOG

| 2025-2026 |



ハイエンドGPUを最大4基搭載可能な AIワークステーション

AIの研究開発、大規模言語処理に最適なワークステーションです。これらのモデルは、機械学習、大規模データ処理、トレーニング、推論などを効率的に行うために設計されており、整合性と動作確認が取れた各種フレームワークをプレインストールしています。

1
CPU
3
GPU

Deep Learning BOX III

インテル® Xeon® W プロセッサ W3500シリーズを採用し、シングルソケットで最大60コア。最新のGPUに対応し最大3基搭載可能。Dual PSU設計で2000W以上の電源を供給するためLLMなどの複雑で大規模なAI開発にもご利用いただけます。

1
CPU
4
GPU

Deep Learning STATION II

インテル® Xeon® W プロセッサ W3500シリーズを採用し、シングルソケットで最大60コア。当社独自のケース設計で最新のGPUを最大4基搭載可能。また、Dual PSU設計により2000W以上の大容量電源を供給するためLLMやロボティクス開発にも適しています。

ローカルLLM・RAG開発に最適な
GPU環境とQ&Aサポート付

Q&A
チケット付

クローズドな環境で機密性の高いデータを用いた開発が可能に

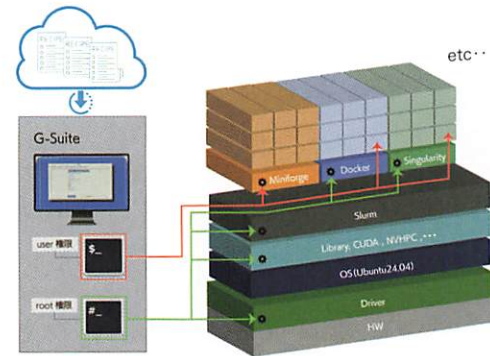
ローカルRAGスターター BOX

ローカルRAGスターター BOXは、RAGとローカルLLMを組み合わせて構築したLLMの開発を、よりお手軽・迅速・安全に実現するGPUワークステーションモデルです。オンプレミス環境で自社固有の文書やデータを活用したRAGを実装することで、機密性の高い自社データを活用した開発が可能になります。

アプリケーションリソース マネージメントツール 「G-Suite」

G-Suite (G-Works LLM Application Suite for Linux ジー・スイート) は、これまで当社が提供してきた「G-Works」の仕組み (Deep Learning の主要なフレームワークを、各世代のGPU に最適化してビルドしたソフトウェア群) を継承し、さらに発展させたOS ネイティブ環境+コンテナ環境のハイブリッドアプリケーションリソースマネージメントツールです。

G-Suite は、インターネット上のWeb サイトから様々なレシピやサポートツール群をダウンロードし、簡易解説書やチュートリアルを含むアプリケーション実行環境を構築することが可能です。最適化されたフレームワークが実装されているためユーザーが面倒な構築作業を行うことなく届いたその日からAIの研究開発を開始することが可能です。



手軽に、高性能に。 「ROBODEV」でAIロボット開発を加速。

— Isaac Simでロボティクスの可能性を広げる —

ROBODEVは、AIロボット開発に携わるお客様が必要な環境構築として「ROBODEVインストーラー」を実装したGPUワークステーションです。ロボットの操作シミュレーションができるNVIDIA Isaac SimをはじめとしたNVIDIAのAIロボット開発プラットフォームが簡単にセットアップできます。セットアップ後は、AIロボット開発導入支援に向けたNVIDIA Isaac Simの使い方をハンズオンサポートするメニューをご用意しており、座学や基本操作をレクチャーします。



最大2GPUまで搭載可能な コストパフォーマンスモデル

GPU演算に加えて、高速演算やシミュレーション、グラフィック作業などCPUパワーを必要とする作業においてコストパフォーマンス性が高いモデルです。最新のNVIDIA GPUの搭載にも対応しています。

1
CPU
2
GPU

GWS-ASTGA-2G

AMD Ryzen™ Threadripper PRO 7000 WX プロシリーズプロセッサ採用によりシングルソケットで最大96コア。最新のGPUに対応し最大2基搭載可能で、CPUパワーを必要とする高速演算等の作業に最適なパフォーマンスを提供します。

1
CPU
2
GPU

GWS-ISAGA-2G

インテル® Xeon® W プロセッサ W3500シリーズを採用し、シングルソケットで最大60コア。最新のGPUに対応し最大2基搭載可能で、GPU演算に加えて、高速演算などCPUパワーを必要とする作業にも満足のいくパフォーマンスを提供します。

1
CPU
2
GPU

GWS-ISASM-2G

インテル® Xeon® Wプロセッサ W3500シリーズを採用し、シングルソケットで最大56コア。最新のGPUに対応し最大2基搭載可能で、複雑な3D CAD、および AI開発とエッジ導入向けの拡張されたプラットフォーム機能も提供します。

大規模言語処理、シミュレーションに適したマルチGPUモデル

最新のインテル® Xeon® スケーラブル・プロセッサやAMD EPYC™ 9005プロセッサを搭載。
GPUはエントリーからハイエンドまで用途に応じて様々なカードが搭載可能なモデルです。



GSV-IEMSM-4U4G

第5世代 インテル® Xeon® スケーラブル・プロセッサを採用し、デュアルソケットで最大128コア。最新のGPUに対応し最大4基搭載可能で、生成AIによる大規模言語モデルのAI開発、HPC研究などにパフォーマンスを発揮します。



GSV-IGRSM-3U8G

第6世代インテル® Xeon® プロセッサを採用しデュアルソケットで最大128コア。最新のGPUに対応し最大8基搭載可能で、大規模AIの学習やロボティクス開発、LLMトレーニングまで要求の高いワークロードに最適です。



GSV-IGRSM-5U8G

第6世代インテル® Xeon® プロセッサを採用しデュアルソケットで最大128コア。最新のGPUに対応し最大8基搭載可能で、冷却性能にも優れているため幅広い分野での高負荷ワークロードに最適です。



GSV-ATRSM-5U8G

第5世代 AMD EPYC™ 9005プロセッサを採用し、デュアルソケットで最大384コア。最新のGPUに対応し最大8基搭載可能で、大規模シミュレーションやロボティクス開発に最適な高速・大容量メモリ構成に最適です。

NVIDIA HGX™ プラットフォームサーバー

NVIDIA HGX™ プラットフォームは、NVIDIAが提供するNVIDIA GPU、NVIDIA NVLink™、NVIDIA ネットワーキング、AI およびHPCソフトウェアスタックのパフォーマンスを最大限引き出せるよう設計されています。

NVIDIA HGX™ AIスーパーコンピューティングプラットフォームに準拠した最新のHGXシリーズのGPUを搭載しています。



HGX-IEMSM-4U4G

第5世代 インテル® Xeon® スケーラブル・プロセッサを採用し、デュアルソケットで最大128コア。GPUはNVLinkで4基のGPUを接続し、小規模なエンタープライズから大規模な統合GPUクラスターまで効率的なスケラビリティを提供します。



HGX-AGESM-8U8G

第5世代 AMD EPYC™ 9005プロセッサを採用し、デュアルソケットで最大384コア。高速フラッシュストレージや広帯域インターコネクトなど拡張性も確保。抜群のコストパフォーマンスでユーザーのAI・HPCの開発を加速します。



HGX-IEMSM-10U8G

第5世代 インテル® Xeon® スケーラブル・プロセッサを採用し、デュアルソケットで最大128コア。HGX™ プラットフォームに準拠した最新のGPUに対応したGPUを最大8基搭載可能で、数兆パラメーターを有する大規模AIモデルの学習やロボティクス開発にもご利用いただけます。

大容量・高信頼性ストレージサーバー / NAS

第4世代インテル® Xeon® スケーラブル・プロセッサ搭載のストレージサーバー、AMD Ryzen™ V1000シリーズのNASモデルです。

エンタープライズクラスの高信頼性のストレージとの組み合わせや、高速ネットワーク接続を提供するNASの利用で安定した運用が可能になります。



FSV-SM2U12B

第4世代インテル® Xeon® スケーラブル・プロセッサを採用し、ストレージベイは前後で合計12スロット装備し様々なRAIDレベル0、1、5、6、10、50、60に対応しています。



FSV-SM4U60B

第4世代インテル® Xeon® スケーラブル・プロセッサを採用し、4Uラックマウント筐体にトップローディングで最大60基のストレージを搭載可能です。また、RAIDレベル0、1、5、6、10、50、60に対応しています。



GNAS-QN2U12B

AMD Ryzen™ V1000シリーズV1500B クアッドコアプロセッサを採用。2つの2.5GbEポートと2つのPCIe Gen 3スロットを備えており、システムのアップタイムを最大化するために冗長電源を搭載しています。



GNAS-QNDS8B

AMD Ryzen™ V1000シリーズV1500B 4コアプロセッサを採用。2つの2.5GbEポートと2つのPCIe Gen 3スロットを備えており、システムのアップタイムを最大化するために冗長電源を搭載しています。

広帯域、低遅延のネットワークスイッチ

企業やデータセンター向けに設計され、信頼性、パフォーマンス、セキュリティを兼ね備えた高性能なネットワークスイッチです。
複雑なネットワーク環境においてもスムーズな通信を実現します。



MQM9700-NS2F

NVIDIA Quantum-2ベースで、NVIDIA高速インターコネクト400Gb/秒のInfiniBandスループットを持ち、AI研究開発向けに高速、超低遅延、スケラブルなネットワーク性能を提供します。



SN2201

データセンター アプリケーション専用に設計された第2世代の NVIDIA Ethernet スイッチです。最大48個の1G/100M/10M Base-T ホストポートにノンブロッキングのスパイン アップリンクで接続します。



SN4600

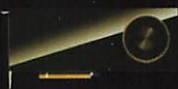
NVIDIA® Spectrum-3をベースとし、クラウド機能と画期的なパフォーマンスを組み合わせたイーサネットスイッチで、最新のスケールアウト分散データセンター アプリケーションをサポートします。



GPU ACCELERATOR GPUアクセラレーター

NVIDIA® GPU ACCELERATOR製品は、最も要求の厳しいHPCやハイパースケールデータセンターのワークロードを高速化します。研究開発における幅広いアプリケーションで直面する課題に対して、従来のCPUよりもはるかに高速に数ペタバイトのデータを解析できるようにします。NVIDIA® H100は、AI、HPC、およびグラフィックスを高速化するために構築された最先端のデータセンター GPUです。

製品 スペック

	NVIDIA RTX 6000 Ada	NVIDIA RTX PRO™ 5000 Blackwell	NVIDIA RTX PRO™ 6000 Blackwell Max-QWorkstation Edition	NVIDIA RTX PRO™ 6000 Blackwell Workstation Edition	NVIDIA RTX PRO™ 6000 Blackwell Server Edition
					
アーキテクチャ	NVIDIA Ada Lovelace Architecture	NVIDIA Blackwell Architecture	NVIDIA Blackwell Architecture	NVIDIA Blackwell Architecture	NVIDIA Blackwell Architecture
メモリ容量	48GB GDDR6	48 GB GDDR7 with ECC	96 GB GDDR7 with ECC	96 GB GDDR7 with ECC	96 GB GDDR7 with ECC
メモリ帯域	最大 960GB/s	1,344 GB/s	1,792 GB/s	1,792 GB/s	1,597 GB/s
CUDA コア数	18,176	12,800	24,064	24,064	24,064
RT コア数	142	100	188	188	188
メモリバス幅	384-bit	384-bit	512-bit	512-bit	512-bit
Tensor コア数	568	400	752	752	752
グラフィックス バス	PCI Express Gen 4 x 16	PCI Express 5.0 x 16	PCI Express 5.0 x 16	PCI Express 5.0 x 16	PCI Express 5.0 x 16
ディスプレイポート	DP 1.4 (4)	DP 2.1 (4)	DP 2.1 (4)	DP 2.1 (4)	DP 2.1 (4)
フォームファクター	112mm(H) x 267mm(L)	112mm(H) x 267mm(L)	112mm(H) x 266mm(L)	137mm(H) x 304mm(L)	112mm(H) x 267mm(L)
最大消費電力	300W	300W	300W	600W	600W

	NVIDIA® H200 NVL	NVIDIA® H200 SXM		NVIDIA® HGX B200	NVIDIA® HGX B300
					
フォーム ファクター	PCIe デュアルスロット空冷	SXM	フォーム ファクター	8x NVIDIA Blackwell SXM	8x NVIDIA Blackwell Ultra SXM
FP64	30 TFLOPS	34 TFLOPS	FP64/FP64 Tensor コア	296 TFLOPS	10 TFLOPS
FP64 Tensor コア	60 TFLOPS	67 TFLOPS	FP32	600 TFLOPS	600 TFLOPS
FP32	60 TFLOPS	67 TFLOPS	TF32 Tensor コア*	18 PFLOPS	18 PFLOPS
TF32 Tensor コア²	835 TFLOPS	989 TFLOPS	FP16/BF16 Tensor コア*	36 PFLOPS	36 PFLOPS
BFLOAT16 Tensor コア²	1,671 TFLOPS	1,979 TFLOPS	FP8/FP6 Tensor コア*	72 PFLOPS	72 PFLOPS
FP16 Tensor コア²	1,671 TFLOPS	1,979 TFLOPS	FP4 Tensor コア**	144 PFLOPS 72 PFLOPS	144 PFLOPS 105 PFLOPS
FP8 Tensor コア²	3,341 TFLOPS	3,958 TFLOPS	INT8 Tensor コア*	72 POPS	2 POPS
INT8 Tensor コア²	3,341 TFLOPS	3,958 TFLOPS	メモリ合計	1.4TB	最大 2.3TB
GPU メモリ	141GB	141GB	合計 NVLink 帯域幅	14.4TB/秒	14.4TB/秒
GPU メモリ帯域幅	4.8TB/秒	4.8TB/秒	NVIDIA NVSwitch™ の特長	NVLink 5 スイッチ	NVLink 5 スイッチ
最大熱設計電力 (TDP)	最大 600W (設定可能)	最大 700W (構成可能)	NVSwitch GPU から GPU への帯域幅	1.8TB/秒	1.8TB/s

NVIDIA のロゴは、NVIDIA Corporation の商標または登録商標です。すべての会社名および製品は、関連各社の商標または登録商標です。機能、価格、提供状況および仕様は予告なしに変更されることがあります。© 2025 GDEP Advance. All rights reserved. 製品画像は実際と異なる場合があります。



株式会社ジーデップ・アドバンス



www.gdep.co.jp

☎ 03-6803-0620 ✉ @GDEPAdvance

〒104-6205 東京都中央区晴海1丁目8-12
晴海アイランド トリトンスクエア オフィスタワーZ棟5階

販売店

大阪大学生協同組合
購買公費センター
〒565-0871 吹田市山田丘2-1 センテラス
TEL 06-6827-6760 大学内線2714