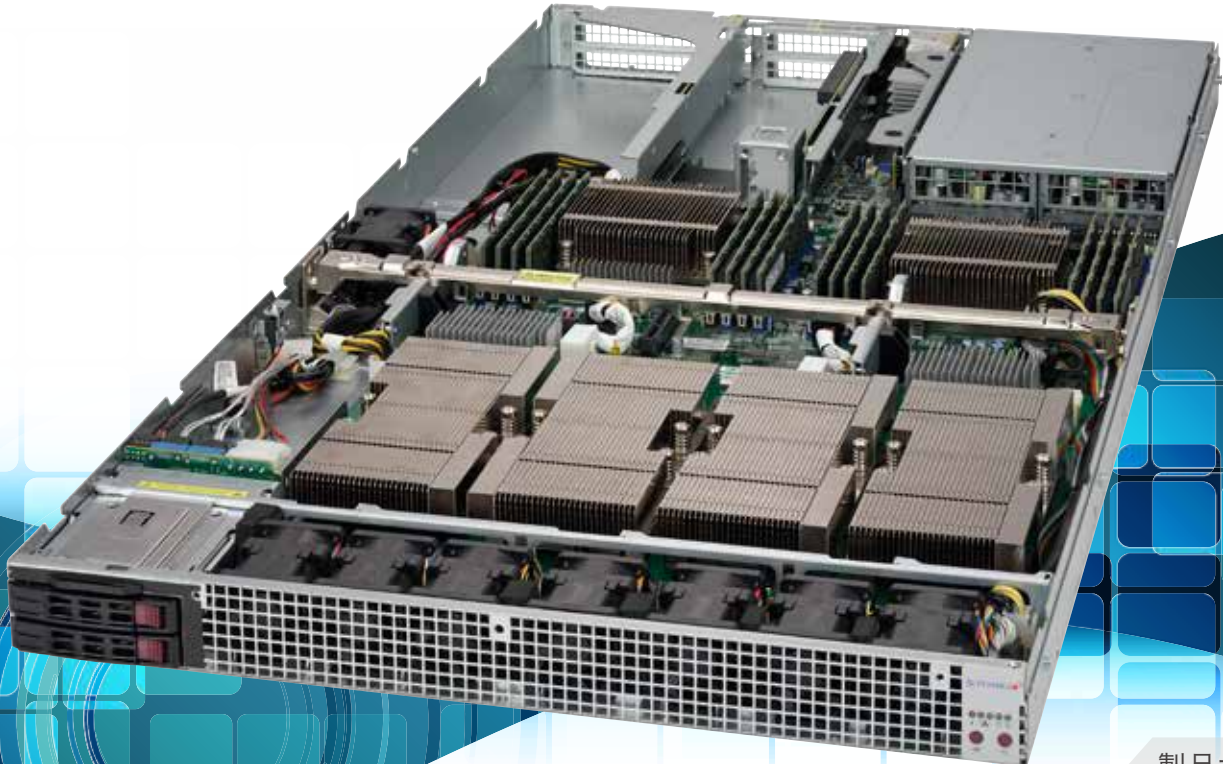


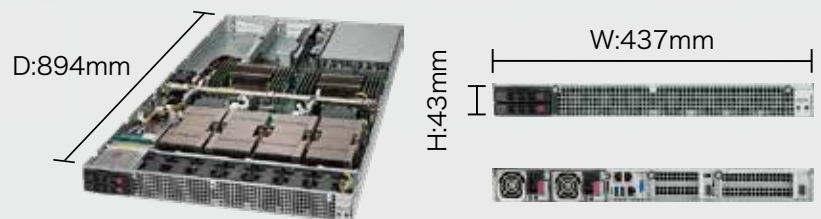
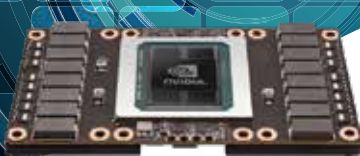
ハイエンドGPUサーバー

HPC5000-XBWGPU4R1S-PCL

Pascal アーキテクチャ採用の新世代GPU「Tesla P100」を4基搭載
HPC、Deep Learningに威力を発揮するハイエンドGPUサーバー



製品サイズ



特長

- 新世代GPU NVIDIA® Tesla® P100 NVLink対応モデルを4基搭載可能
- インテル® Xeon® プロセッサー E5-2600 v4 ファミリー対応
- 最大2CPU (44コア)、最大1TBメモリ搭載可能
- 安定的な運用を確保する冗長化電源を搭載 (80PLUS TITANIUM認証取得)
- IPMI2.0が 高度な遠隔監視、操作を実現
- 省スペースな1Uラックマウント筐体
- 深層学習に必要な主なソフトウェアのインストールサービスが付属



製品仕様

HPC5000-XBWGPU4R1S-PCLは、新アーキテクチャ「Pascal」をベースとした最新の数値演算アクセラレータ NVIDIA® Tesla® P100 NVLink対応モデルを2基または4基搭載することができます。

製品名	Tesla P100 for NVLink-enabled Servers
アーキテクチャ	Pascal
CUDA コア	3584
コアクロック	1.328GHz（GPU Boost 時最大 1.480GHz）
低精度浮動小数点演算性能	4.76TFLOPS（GPU Boost 時 5.30TFLOPS）
単精度浮動小数点演算性能	9.52TFLOPS（GPU Boost 時 10.61TFLOPS）
半精度浮動小数点演算性能	19.04TFLOPS（GPU Boost 時 21.22TFLOPS）
NVLink 帯域幅	160GB/s（双方向）※
PCIe x16 帯域幅	32GB/s（双方向）
メモリ容量	16GB
メモリ帯域幅	732GB/s
消費電力	300W

※NVIDIA® Tesla® P100 for NVLink-enabled Servers (GP100) は、高速インターコネクト「NVLink」を4リンク備えています。NVLinkによるGPU間の接続帯域幅は1リンクあたり双方向40GB/s。4リンク合計で双方向160GB/sとなります。

HPC50000-XBWGPU4R1S-PCLは、14nm世代のインテル® Xeon® プロセッサー E5-2600 v4 ファミリーを2CPU搭載しています。最上位モデルのE5-2699 v4 (22コア, 2.2GHz) を選択することで、最大44コアまで実装することができます。

HPC5000-XBWGPU4R1S-PCLは、64GBメモリモジュール (DDR4 LRDIMM-2400 Registered ECC) を16本のメモリスロットに搭載する事で最大1TBのメモリ容量を確保します。メモリ性能を必要とする大規模な計算でパフォーマンスを発揮します。

HPC5000-XBWGPU4R1S-PCLは、2.5型 HDD/SSDを2台まで搭載可能です。

※標準構成では240GB SSDを2台搭載しています。

HPC5000-XBWGPU4R1S-PCLは、80PLUSで最上位ランクの80PLUS TITANIUM認証を取得した高効率な電源を搭載しています。80PLUS認証とは、交流から直流への変換効率を保証するものです。80PLUS TITANIUM認証は、負荷率10%/20%/50%/100%でそれぞれ90%/92%/94%/90%という高い変換効率基準をクリアしたものに与えられます。

HPC5000-XBWGPU4R1S-PCLは、100Vから240Vに対応した2000W電源ユニットを2個搭載し、一方の電源ユニットに障害が発生した場合でもサーバーの運転を継続するための電力を充分に供給できる冗長性を持っています。これにより万が一の電源ユニット障害によるダウンタイムを最小限に抑えることが出来ます。

標準搭載されたIPMI2.0機能は専用のLANポートを備え、リモートによる温度、電力、ファンの動作、CPUエラー、メモリーエラーの監視を可能にします。また電源のオンオフ、コンソール操作を遠隔から行うことができます。これらの機能によりシステムの信頼性、可用性を高め、ダウンタイムとメンテナンス費用を圧縮することを可能にします。

本製品には、深層学習に必要な主なソフトウェアのインストールサービス※が付属します。

OS: Ubuntu 16.04 LTS または Ubuntu 14.04 LTS
CUDA Toolkit: CUDAを拡張したGPUユニバーサルライブラリー、ドライバ、ツールなどが含む統合開発環境
cuDNN: Deep Neural Network (DNN) 用のCUDAライブラリー
Caffe/**PyCaffe**: オープンソースの Deep Learning Framework および Python で使っための PyCaffe
Torch: 古くからあるオープンソースの Deep Learning Framework
Caffelet: Preferred Networksが開発したオープンソースの Deep Learning Framework
TensorFlow: GoogleのAI開発環境を一般向けにカスタイズしたオープンソースの Deep Learning Framework
DLGITS: Deep Neural Network の構築がすばやく簡単に行えるソフトウェア
NCCL: マルチGPU集合通信ライブラリー

多くのアプリケーションが續々とCUDAに対応しています。HPCシステムズのHPC5000-XBWGPU4R1S-PCLなら、CUDA化されたアプリケーションの活用に最適です。

製品名	HPC5000-XBWGPU4R1S-PCL
OS	【HPC向け】CentOS, Red Hat Enterprise Linux 【Deep Learning向け】Ubuntu ※Windowsを希望される場合は、別途ご相談ください。
プロセッサー	インテル® Xeon® プロセッサー E5-2600 v4 ファミリー ※CPUモデル (プロセッサーナンバー) によって設置環境に制限があります。詳細はお問い合わせください。 ※E5-2687W v4 / 2667 v4 / 2643 v4 / 2637 v4 / 2623 v4 は搭載できません。
プロセッサー搭載数	2CPU (44コア)
プロセッサー冷却方式	空冷式
チップセット	インテル® C612
メモリ	1TB (64GB DDR4 LRDIMM-2400 Registered ECC ×16) 512GB (32GB DDR4 2400 Registered ECC ×16) 256GB (16GB DDR4 2400 Registered ECC ×16) 128GB (16GB DDR4 2400 Registered ECC ×8) 64GB (8GB DDR4 2400 Registered ECC ×8)
メモリスロット	16DIMMsロット/DDR4 LRDIMM-2400 Registered ECC (64GB), DDR4 2400 Registered ECC (8,16,32GB)
GPU	NVIDIA® Tesla® V100 for NVLink-enabled Servers 32GB NVIDIA® Tesla® V100 for NVLink-enabled Servers 16GB
GPU搭載数	4基
ハードディスクドライブ	標準 : 240GB SSD (2.5型, SATA) ×2 ※HDD/SSD (2.5型, SATA) を最大2台搭載可能 ※RAIDアレイコントローラー (オプション) 増設時, SAS HDD使用可能
光学ドライブ	なし
グラフィックス	Aspeed AST2400
インターフェイス per node	VGA [D-sub15ピン] (背面) ×1 USB3.0 (背面) ×2 ネットワーク [GbEポート] (背面) ×2 IPMI2.0ポート [RJ45] (背面) ×1
拡張スロット	PCI-Express 3.0 (x16) ×3, PCI-Express 3.0 (x8) ×1 (Low Profile)
電源ユニット	2000W 冗長化電源 (80PLUS TITANIUM認証取得)
ACケーブル	200V用ACケーブルを2本添付 / IEC320-C13 ⇒ IEC320-C14
ACコネクタタイプ	IEC 320-C14
最大消費電力	1826W
筐体タイプ	ラックマウントタイプ (1U)
サイズ (縦幅×横幅×奥行)	43mm × 437mm × 894mm
重量	16 kg
付属品	200V用ACケーブル ×2 USBキーボード (日本語または英語) ×1 USB光学式スクロールマウス ×1 取扱説明書 保証書
オプション	RAIDアレイコントローラー 2.5型SSD (フラッシュメモリドライブ) InfiniBand HCA 各種ディスプレイ
保証	3年間センドバック保守

販売店	

