



DellのAIポートフォリオでAI成功の道筋を見つける

DellのAIポートフォリオとSupermicroの類似製品の比較

人工知能(AI)は新たなフロンティアであり、あらゆる産業のビジネスオペレーションの形を変えようとしています。あらゆる組織が、ビジネスオペレーションの向上にAIを活用する方法を模索していますが、ここで忘れてはならないのは、AIを導入してそのメリットを手にすることが一夜にして達成できるわけではないことです。ビジネスはどれも同じではないため、自社のデータと自社のビジネスゴールを見直して、AIをどのようにそのデータとともに使用すれば望むとりの成果が得られるかを考える必要があります。Dellのような、包括的なAIポートフォリオ（計画、データ準備、適切なハードウェア選択、AIモデル設計、概念実証テスト、リファレンスアーキテクチャ、エンドツーエンドのサポートなど）を提供する企業の協力を得ることが、AIプロジェクトの成功につながります。

市場には無数の選択肢があるため、意思決定のすべてを支援できるパートナーが見つかるかどうか、AI導入の成功と手痛い失敗の分かれ道といえます。本レポートではDellとSupermicroのAIポートフォリオを精査しますが、そのねらいは顧客のAI導入の過程全体を通してDellが持つ優位性を知っていただくことです。最初に焦点を当てるのはサーバーとコンピューティングのオプションですが、この領域では両社が多数の多様な製品を提供しています。次に、Dellがハードウェアだけでなく教育、計画サービス、パートナーエコシステムなどを求める企業にどのように応えているかを考察します。

AIワークロードのためのサーバーとパフォーマンスの成績

サーバーは、AIワークロード実行の基礎となるコンピューティング インフラストラクチャであり、ワークロードのサイズやタイプに応じてCPUまたはGPU、またはその両方をコンピューティング リソースとして使用できます。大規模な、または大量のリソースを必要とするワークロード、たとえばHPCやAIに対しては、GPUを使用すると最高のパフォーマンスを達成できます。GPUのフォーム ファクターはユニバーサルPCIeやOpen Compute Project Accelerator Module (OAM)などさまざまなものがあり、なかでもNVIDIA独自のSXMアーキテクチャのパフォーマンスは現在最高レベルです。¹ また、冷却アーキテクチャや電力効率といったサーバー設計や大容量メモリーもパフォーマンスに影響します。多くのデータセンターでは依然として空冷が使用されているため、AIワークロードのためのサーバーは可能な限り効果的に空気で冷却するように作られている必要があります。ここでは、Dell PowerEdgeサーバー 製品群をコンポーネントと冷却オプションなどの観点から、一般公開されているMLCommons® MLPerf®スコアとともに紹介します。

テスト結果

MLPerf®は、AIのトレーニングと推論性能を評価するベンチマーク スイートです。組織がMLPerf®の公式結果を公開する場合、その結果はベンチマーク開発者MLCommons®によって設定された所定の条件に準拠している必要があります。² これらのコンプライアンス ガイドラインは、性能比較を容易にする基準を提供するものです。MLPerf®では、Datacenter、Edge、Mobile、Tinyデータセットを使用して推論テストを実施し、AIスコアとテスト中の電力消費ワット数を報告しています。推論ベンチマーク スイートには、多数の一般的なAI、ML、DLモデルのテストが含まれています（表1を参照）。

表1： MLPerf®に含まれるAI、ML、DLのモデルと、それぞれのテストおよび標準的な用途。出典：Principled Technologies

一般的なAIモデル	標準的な用途
ResNet	画像分類モデル。コンピューターに様々な画像を学習、記憶、識別させる場合に役立つ。ユース ケースとしては、医療画像、ソーシャル メディア コンテンツのモデレーション、顔認識など
RetinaNet	オブジェクト検出の一種であり、ResNetよりもさらに複雑なものを扱うことができる。画像やビデオ フレームに含まれているオブジェクトの識別と検出に役立ち、それらのオブジェクトを重要度によって分類できる。自動運転、車両自動アシスト技術、監視、顔認識などに使用される
3D-UNet	医療画像のセグメント化に特化
RNN-T	音声認識。言語の自動翻訳などのユース ケースで使用される
BERT	自然言語処理。テキストの要約、言語翻訳、タスクの自動完了などのユース ケースで使用される
DLRM-v2-99.9	推奨モデル。ターゲットを絞った広告やパーソナライズされた製品推奨などのユース ケースで使用される
GPTJ-99と99.9	自然言語処理用LLM。テキスト生成に優れており、チャットボットやチャットベースのAIツールといったユース ケースで使用される



MLPerfについて

MLPerf®の結果には、AIモデル自体のパラメーターとは別にいくつかの追加パラメーターが含まれています。これらを利用して、大量のデータを1つのグラフまたは表内で解析できます。パラメーターについて簡単に説明します。

- 99.0と99.9：これらの数値は、モデルがどれだけの精度でトレーニングされたかを表します。より正確な出力を必要とする場合、モデルはより複雑になり、データの処理時間も長くなります。
- オフライン サンプル/秒：このベンチマークモードでは、すべてのクエリーがテスト開始時に送信されます。システムにすでに存在するデータのシミュレーションです。
- サーバー クエリー/秒：このベンチマークモードでは、クエリーがテスト期間全体にわたって送信されます。データのライブストリーム分析のシミュレーションです。

MLCommons®およびMLPerf®の結果の詳細については、

<https://mlcommons.org/benchmarks/inference-datacenter/>を参照してください。

本レポートで取り上げている結果は、MLCommons®のWebサイトで2023年11月に公開されたMLPerf® v3.1 Inference Datacenterの結果です。³これらの結果にはテクノロジー メーカーやクラウド サービス プロバイダーから提出された結果が含まれており、幅広い構成を網羅しています。一般公開されているSupermicroからの提出データと比較して、Dell PowerEdgeサーバーは同等の結果を達成しています。表2にサーバーの詳細を示します。

表2：2023年11月に公開されたMLCommons® MLPerf® 3.1でのDellおよびSupermicroのサーバーの結果出典：Principled Technologies

送信者	サーバーのモデル	GPUの数とモデル	説明
Dell ⁴	PowerEdge XE9680	8x NVIDIA H100 SXM	大規模ワークロード（大規模言語モデルなど）を使用したAIトレーニングおよび推論向け
	PowerEdge XE9640	4x NVIDIA H100 SXM	高密度の水冷データセンターにおける大規模AIモデルのトレーニング向け
	PowerEdge XE8640	4x NVIDIA H100 SXM	従来型のAIトレーニング、HPC、データ分析アプリケーションの駆動向け（4Uフォーム ファクター、空冷データセンター向け）
Supermicro ⁵	AS-8125GS-TNHR	8x NVIDIA H100 SXM	大規模AIトレーニングおよびHPCワークロード向け（AMDプロセッサ搭載）
	SYS-821GE-TNHR	8x NVIDIA H100 SXM	大規模AIトレーニングおよびHPCワークロード向け（インテル プロセッサ搭載）
	SYS-421GU-TNHR	4x NVIDIA H100 SXM	HPCおよびAIのワークロードを柔軟にサポートするモジュラー型設計

DellとSupermicroの両社から同一のGPU構成の結果が提出済みのため、パフォーマンスの比較は単純明快です。図1と図2に示すように、両ベンダーの構成に共通部分が多い場合は、結果はほぼ同等です。他の構成では、図3と図4に示すように、4 GPUテストでのgptj-99.9モデルの成績はDellがSupermicroを上回っています。なお、Dellはすべてのモデルについて3サーバーすべての結果を提出済みですが、Supermicroはそうではありません。ここでは、両社の結果が提出済みであるモデルについてのみ比較します。Dellのすべての結果については、MLCommons® MLPerf®の結果にアクセスしてください。

8 GPUサーバーの結果

Dell PowerEdge XE9680は、AIアクセラレーション用のNVIDIA H100 SXM5 GPUを最大8基、第4世代インテル® Xeon® スケーラブル・プロセッサ・ファミリーを最大2基サポートします。PowerEdge XE製品ファミリーはモジュラー型アーキテクチャを採用しており、SXM4またはSXM5 NVIDIA GPUまたはOpen Compute Project Accelerator Module (OAM) GPUアセンブリーがサポートされるため、標準のPCIe GPUと比べてパフォーマンスを大幅に高めることができます。Dell PowerEdge XE9680には、AMD Instinct™ MI300Xアクセラレーターも採用されています。⁶ ラックスペースがわずか6UのPowerEdge XE9680は、コンパクトな8ウェイNVIDIA H100 SXM5サーバーです。

対照的に、8 GPUのSupermicroのサーバーは、必要なラックスペースが33%多くなりますが(8U)、NVIDIA GPU用のSXMフォームファクターも提供しています。このようにサイズが異なるため、Dell PowerEdge XE9680サーバーは1ラックに7台を収容できるのに対し、Supermicroのサーバーは5台のみとなります。図1と図2では、Dell PowerEdge XE9680の結果を、Supermicro 8 GPUサーバーの2つの構成（インテル プロセッサ搭載のSYS-821GE-TNHRとAMDプロセッサ搭載のAS-8125GS-TNHR）の結果と比較しています。なお、SupermicroからSYS-821GE-TNHRでのRNN-Tの結果が提出されていないため、このモデルは図1のグラフから除外されています。

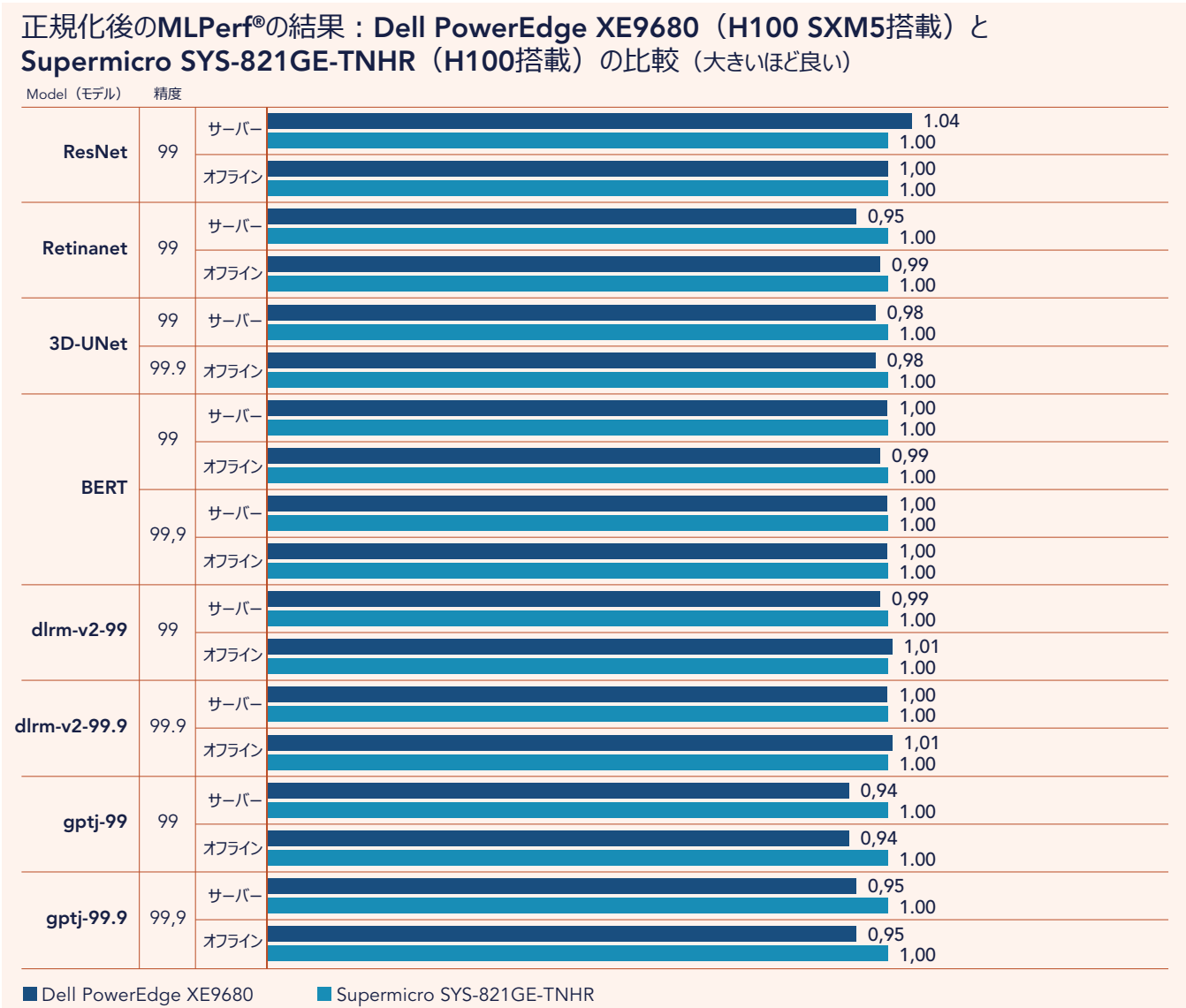


図1：一般公開されているMLPerf®によるDell PowerEdge XE9680とSupermicro SYS-821GE-TNHRの結果（2023年11月29日時点）。どちらのシステムもSXMフォームファクターのNVIDIA H100 GPUを装備している。出典：Principled Technologies。MLCommons®からのデータを使用。^{7,8}



正規化後のMLPerf®の結果：Dell PowerEdge XE9680（H100 SXM5搭載）と Supermicro AS-8125GS-TNHR（H100 SXM5搭載）の比較（大きいほど良い）

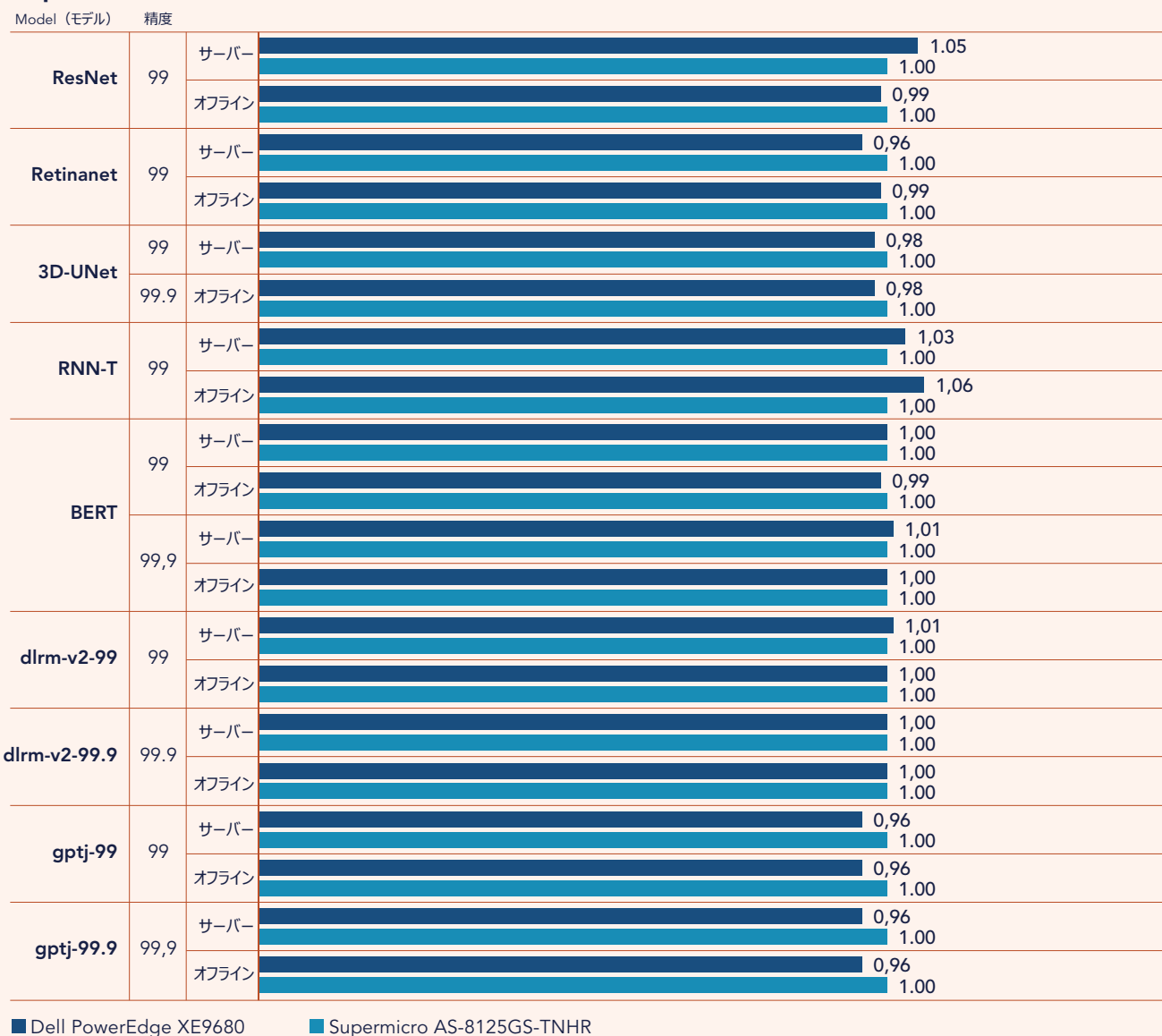


図2：一般公開されているMLPerf®によるDell PowerEdge XE9680とSupermicro AS-8125GS-TNHRの結果（2023年11月29日時点）。どちらのシステムもSXMフォームファクターのNVIDIA H100 GPUを装備している。出典：Principled Technologies。MLCommons®からのデータを使用。^{9,10}

これらの結果が示すように、Dell PowerEdge XE9680を選択すると、多数のAI推論ワークロードでの同等のパフォーマンスを、データセンターの占有スペースを抑えて達成できます。

4 GPUサーバーの結果

データセンターの電力消費またはスペースの最小化が最重要事項であるときは、2UのDell PowerEdge XE9640が解決策になる可能性があります。NVIDIA H100 SXM GPUを最大4基搭載するPowerEdge XE9640の特長は、PowerEdge XE9680の半分のGPU計算能力を3分の1のスペースで達成できることです。¹¹ 高密度設計のDell PowerEdge XE9640サーバーはDell Smart Coolingテクノロジーを採用しており、CPUとGPUに対する直接水冷など、多様なサーマルテクノロジーを実現しています。¹² PowerEdge XE9640の2Uシャーシは、大型ファンやヒートシンクなど向上したエアフローメカニズムに対応しており、PCIeカードやメモリーといったその他の重要コンポーネントの冷却に役立ちます。¹³

Supermicroからは、これよりも前にSYS-220GQ-TNAR+サーバーが発表されており、NVIDIA A100 HGX GPUを4基搭載した2Uフォームファクターが特長ですが、PowerEdge XE9640に匹敵する、新しいH100 HGX GPUを4基搭載したSupermicroの2Uサーバー製品は見つかりませんでした。¹⁴ HGX H100 NVIDIA GPUを搭載する4 GPUサーバーとしてSupermicroからMLPerf®に提出されているのはSYS-421GU-TNXRですが、これは4Uサーバーです。前述したように、Supermicroから提出済みのMLPerf®3.1の結果はSYS-421GU-TNXRでのgptj-99.9 AIモデルに対するもののみであるため、他のモデルについてはPowerEdge XE9640との比較はできません。しかし、一般公開されている結果では、PowerEdge XE9640はオフラインテストでの成績がSupermicroサーバーを上回り、最大1.37倍のスコアを達成しています（図3を参照）。

正規化後のMLPerf®の結果：Dell PowerEdge XE9640（H100 SXM5搭載）とSupermicro SYS-421GU-TNXR（H100 SXM5搭載）の比較（大きいほど良い）

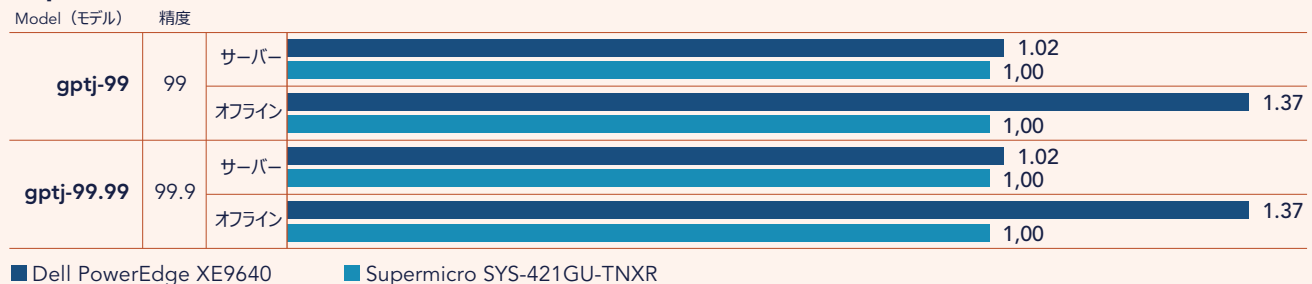


図3：一般公開されているMLPerf®によるDell PowerEdge XE9640とSupermicro SYS-421GU-TNXRの結果（2023年11月29日時点）。どちらのシステムもSXMフォームファクターのNVIDIA H100 GPUを装備している。出典：Principled Technologies, MLCommons®からのデータを使用。^{15,16}

また、Supermicro SYS-421GU-TNXRサーバーの結果をDell PowerEdge XE8640と比較することもできます。この4U 4 GPUサーバーもNVIDIA H100 HGX GPUをサポートしています。PowerEdge XE9640よりも大型ですが、PowerEdge XE8640は直接水冷を必要としないため、水冷を利用できないデータセンターにとって密度と冷却テクノロジーのバランスが取れた製品となっています。PowerEdge XE8640では、プロセッサには空冷、GPUには液体アシスト空冷ラジエーターが使用されており、施設からラックへの水供給を必要としません。¹⁷ Dell PowerEdge XE8640の特長は、最新の第4世代インテル Xeon スケーラブル プロセッサと最大4 TBのメモリーであり、AIとデータ分析では一般的である大規模データセットと複雑な計算を扱うことができます。¹⁸ 4UフォームファクターのPowerEdge XE8640は、密度とGPU能力のどちらもSupermicro SYS-421GU-TNXRと同等です。しかし、PowerEdge XE9640と同様に、Dell PowerEdge XE8640はオフラインテストでのgptj-99のスコアがSupermicroのサーバーを上回っています（図4）。

正規化後のMLPerf®の結果：Dell PowerEdge XE8640（H100 SXM5搭載）と Supermicro SYS-421GU-TNXR（H100 SXM5搭載）の比較（大きいほど良い）

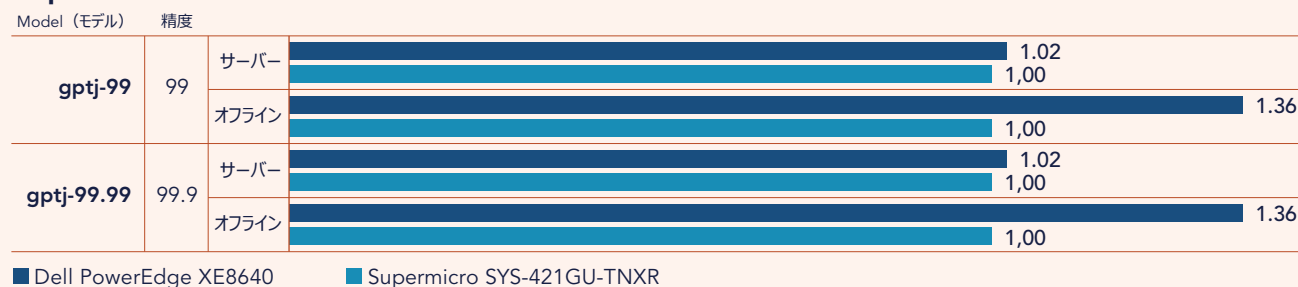


図4：一般公開されているMLPerf®によるDell PowerEdge XE8640とSupermicro SYS-421GU-TNXRの結果（2023年11月29日時点）。どちらのシステムもSXMフォームファクターのNVIDIA H100 GPUを装備している。出典：Principled Technologies, MLCommons®からのデータを使用。^{19,20}

これまで見てきたように、SupermicroとDellのGPUベースサーバー製品群のMLPerf®パフォーマンス結果は全体的に同等ですが、あるAIモデルについてはDell PowerEdge XE8640およびPowerEdge XE9640サーバーに顕著な優位性があります。

パフォーマンスはAI導入の過程における考慮事項の1つにすぎないため、本レポートでは、DellおよびSupermicroが提供するその他のAI関連製品群、つまりワークステーション、ストレージ、ネットワーキングやサービス、サポート、教育などにも注目しました。その結果判明したのは、DellのAIポートフォリオはSupermicroの製品群よりも広範であり、コンピューティングのパフォーマンスの他にも企業が直面する問題の多くに対するソリューションが提供されていることです。

クライアントワークステーションとストレージの製品群

追加のコンピューティングオプションとしてのワークステーション

AIのユースケースの中には、別のコンピューティングアプローチを必要とするものがあります。AI対応ハードウェアへのアクセスを必要とする人の誰もがデータセンターに縛られたままでいられるとは限りません。研究者の研究室には、サーバーを収容したラックを置くスペースがないこともあります。また、外で活動するときに、大型で重いデスクトップシステムを持ち歩くのは現実的ではありません。そのようなときにワークステーションの出番となります。DellのAIポートフォリオにはAI対応のさまざまなPrecisionワークステーションが含まれており、さまざまなニーズを満たせるようにタワー、モバイル、ラック構成が用意されています。²¹ 対照的に、Supermicroからは、外出先でも簡単にアクセスするためのモバイルワークステーションは提供されていません。同社のGPUワークステーション製品群にはさまざまなタワー構成の製品があり、Supermicroによるとそのいくつかはユーザーがラックマウントできます。²² 当社の調査ではラックマウント可能ワークステーションが2種類見つかりましたが、どちらも最新ではなく、搭載しているのはNVIDIA A100 GPUのみであり、すでに販売を終了している可能性があります。²³ 企業が、導入するGPUワークステーションに種類と可搬性に関する柔軟性を求めている場合は、そのニーズを満たすという点でDell AIポートフォリオのほうが優れています。

ストレージに関する考慮事項

ストレージも、AIワークロードを実行するときの重要性はおそらくコンピューティングと同等です。データが多いほどAIモデルの精度は向上しますが、多くのデータセンターでは大量のデータセットの保存と管理は容易ではありません。また、一般的にモデルのトレーニングに使用されるのは非構造化データであるため、AI対応ストレージシステムはさまざまな種類のデータを簡単に処理する必要があります。²⁴ AI、ML、DLのデータセットに対応する容量と拡張性を提供するため、DellはPowerScale™シリーズをファイルストレージおよびElastic Cloud Storage (ECS)用として、またはオブジェクトストレージ用のソフトウェア定義オブジェクトScale用として提供しています。

Dell PowerScaleオールフラッシュNASポートフォリオでは、容量オプションとしてノードあたりのraw容量が3.84 TBから最大720 TBまで用意されており、クラスター化されたオールフラッシュ容量は186 PBに達します。PowerScaleの柔軟性と拡張性は、様々な顧客およびAIのユースケースに対応できます。²⁵ オールフラッシュPowerScaleの3モデル（F200、F600、F900）すべてにインラインデータ圧縮と重複排除機能があり、ストレージ効率が向上しています。²⁶ PowerScaleストレージの各モデルではDell OneFS™ファイルシステムが使用されており、高度な階層化ポリシーが活用されているため、アクセス頻度が最も高いデータが確実に、最高パフォーマンスのストレージ層に配置されます。²⁷ DellはOneFSソフトウェアをAmazon Web Services (AWS)マーケットプレイスでも、APEX File Storage for AWSとして販売しています。OneFSをAWSのコンピューティング インスタンスとともに利用すると、オンプレミスのOneFSアレイと同じ機能が使用できるようになり、ユーザー エクスペリエンスを統一できます。²⁸

Supermicroのストレージ製品群は、ストレージ サーバーで構成され、さまざまなサイズと密度のラック マウント型ストレージ高密度サーバーが提供されています。²⁹ Supermicroのファイル ストレージを利用するには、サードパーティーからのさまざまなソフトウェア デファインド型ストレージ製品群（WekaIO、Scality RING、OSNEXUSなど）から選ぶ必要があります。³⁰ Scality RINGとOSNEXUSは、それぞれのプラットフォームの説明としてファイル ストレージ オプションが含まれていますが、基本的なファイル ストレージを求めている場合はWekaIOが第一の選択肢になるといえます。Supermicroは幅広いユースケースをカバーするさまざまなリファレンス アーキテクチャを用意していますが、利用するにはWekaIOソフトウェアのサブスクリプションまたはライセンスが必要であり、その結果としてソリューションの全体的なコストが増加する可能性があります。³¹

Dellのオブジェクトストレージ オプションにはDell ECS Enterprise Object Storageが含まれていますが、これは「非構造化データをパブリッククラウド規模で保存するという目的に特化」しています。³² Amazon S3オブジェクトストレージとの組み込みの互換性でハイブリッド クラウド機能を提供するほか、ECSストレージ ノードはラックあたり最大14PBの容量を提供します。³³ ファイル ストレージと同様に、Supermicroのオブジェクトストレージ製品群を利用するにはサードパーティー構成が必要です。OSNEXUSプラットフォームは、ファイル、ブロック、オブジェクトストレージを統合したストレージ プラットフォームであり、Scality RINGソリューションはファイルとオブジェクトストレージを組み合わせたものですが、どちらもサードパーティー ベンダーのライセンスが必要です。^{34,35} オブジェクトストレージだけが必要な場合は、SupermicroのQuantum ActiveScale対応ソリューションを購入すると、Quantumソフトウェア サブスクリプションを使用するプライベートクラウド オブジェクトストレージが実現します。³⁶ 本レポートの執筆時点では、Supermicroから提供されている柔軟な使用/運用コスト オプションは見つかりませんでした。

Supermicroのストレージ製品群を利用するにはサードパーティーのソフトウェア ベンダーとの契約が必要になるため、ライセンスまたはサブスクリプションのコストが追加で発生することが多く、サポートやトラブルシューティングなどが簡単ではないことがあります。DellのAIポートフォリオでは、Dellのストレージを購入した顧客が、そのストレージ ソリューションのあらゆる面についてのサービスとサポートを単一の信頼できるソリューションで受けることができます。

Dell APEXについて

ファイル ストレージを「アズ ア サービス」として使用することを望む顧客向けに、DellはAPEX Data Storage Servicesを用意しており、これにはファイル、ブロック、バックアップ ストレージが含まれています。Dell APEX Consoleを使用して、新しいサブスクリプションの注文や、ストレージ容量の調整と監視などを行うことができます。Dellによると、このソリューションによって「クラウド エクスペリエンスの容易さと俊敏性に加えて、自分のアプリケーションとデータに対する詳細な制御」が得られます。³⁷

Dell APEXの詳細については、<https://www.dell.com/en-us/dt/apex/storage/data-storage-services/index.htm>をご覧ください。



ネットワーキング オプション

ネットワーキングも、AIインフラストラクチャの重要なコンポーネントです。多くのAIワークロードが大規模なサーバー クラスタで実行され、サーバー間およびサーバーとストレージの間で絶えず通信が必要になるため、AIワークロードにはボトルネックを回避するための堅牢なネットワーキングが必要になります。ネットワーキングがAIワークロードに対して十分でない場合は、トレーニングと推論に要する時間が増加し、その結果としてデータ処理が遅くなり、インサイトを得るための時間が長くなります。Dellは、PowerSwitch Data Centerトップオブブラック(ToR)スイッチとPowerEdge MX I/OモジュールをEthernetおよびファブリック ネットワーク向けに提供しています。³⁸ PowerSwitchの製品群は、1GbEから400GbEまであり、幅広いニーズを満たすことができます。さらに、Dell PowerSwitch Zシリーズ スwitchは100GbEおよび400GbEの接続が可能であり、リーフ/スパイン ファブリックに合わせて最適化されています。³⁹

Supermicroもスイッチを提供しており、そのEthernetポートは最大400GbEでToRおよびその他のアプリケーション（データセンター スパインとリーフなど）に適しています。⁴⁰ しかし、Dellのネットワーキング サービスには、使いやすさと柔軟性の点でSupermicroにはないさまざまなメリットがあります。Dell Fabric Design Centerなどのサービスは、ネットワーキングのミスマッチ、ギャップ、非効率性の回避に役立つため、自動化を使用するネットワーク ファブリックの計画と導入がしやすくなります。⁴¹ 特定の環境（VMware VxRail、VMware ESXi、Dell PowerStore構成など）向けに、DellはSmartFabric Servicesを提供しており、ソフトウェアデファインド インフラストラクチャの導入とライフサイクル管理が可能になります。PowerStoreとSmartFabric Servicesを組み合わせると、「LAN接続タスクの最大99%を、プラグアンドプレイのファブリックを使用して」自動化できます。⁴² このような、ネットワーキングの設計と実装のための自動化およびガイダンスなどのサービスによって、顧客はサポートを受けながらAI導入プロセスを進めていくことができます。

サービス、トレーニング、その他

AIの導入におけるハードウェア以外の課題の筆頭は、社内に戦略、計画、データ準備、管理に関するエキスパートを持つことの必要性です。AIワークロードの管理と保守を行うには、固有の知識セットが必要であり、これは従来型のハードウェアに関する専門知識に加えて機械学習の運用とデータサイエンスも含まれます。また、AI戦略の設計と実行を担当する人々は、会社独自の経営目標も十分に理解している必要があります。新しいAIワークロードがその目標を確実に達成できるようにするためです。⁴³

別の大きなハードルとして、AIを既存の業務運営システムにシームレスに統合することが考えられます。この統合を行うには、新しいAIテクノロジーと現行のビジネス プロセスとの間の戦略的な整合作業が必要です。AIが導入されても、確立済みのワークフローに影響が及ぶことがないようにするためです。パートナーとして、Dellのようにさまざまな最適化された検証済みソリューション リファレンス アーキテクチャ、トレーニング コース、管理オプション、大規模なパートナー エコシステムを持つ企業を選ぶと、AI導入の過程が容易になります。

Professional Services for AI

教育と計画に関して、DellはAIに特化したさまざまなサービスを提供しています。⁴⁴ AIの導入をサポートするDellのサービスには、コンサルティング、データ準備、展開、サポート、教育が含まれており、それぞれがAI導入における特定の面をターゲットとしています。Dell Advisory Services for Generative AIは、顧客のロードマップ作りを支援するものであり、ユースケースを特定して企業のプロセスの合理化を手助けします。⁴⁵ 同様に、Adoption Services for Generative AIでは、Dellのプロフェッショナルとのワークショップを開催して顧客のニーズと固有の課題をレビューすることにより、顧客のビジネスのための事前トレーニング済みモデルを提案するとともに、顧客のITスタッフをトレーニングする知識伝達セッションも実施します。⁴⁶ また、DellはImplementation, Scale, and Managed Services for Generative AIとしてさまざまなレベルのサポートとトレーニングを提供しており、最大ではフルマネージドのAIインフラストラクチャを提供できるため、顧客のITスタッフはハードウェアの管理をDellに任せ、モデルとデータに専念できるようになります。⁴⁷ ProSupportのサービスによって、最適なシステム パフォーマンスが保証され、継続的なAI運用に必要なハードウェアとソフトウェアのサポートが提供され、技術的な問題が解決されます。⁴⁸

教育サービスは、AI活用に必要なスキルと知識を促進するのに不可欠です。Dellのトレーニング サービスには、データサイエンス分野の包括的なトレーニング プログラムや、高度な分析認定、特定のAIテクノロジー（たとえば機械学習）についてのワークショップが含まれています。⁴⁹

対照的に、Supermicroのサービスはトラブルシューティング、マニュアル、返品保証(RMA)、製品保証に限定されています。⁵⁰ 設計、実装、管理、教育に関するサービスは、SupermicroのAIポートフォリオの中で見つけることができませんでした。AI導入の複雑さを乗り越えるにあたってトレーニング パートナーを探している企業が、この2社のうちDellを選ぶのは明らかです。

AIワークロードに関するサードパーティーとのパートナーシップ

デル・テクノロジーとNVIDIAは共同でDell Validated Designsを提供していますが、そのねらいはビジネスに利用される生成AIのための包括的なソリューションを提供することです。このプロジェクトによって作られる、拡張可能でハイパフォーマンスのインフラストラクチャはDellとNVIDIAのテクノロジーおよびソフトウェアに加えてAIモデル フレームワークをベースにしており、企業はこれを利用してカスタムAIモデルを構築し、実行できます。このソリューションを利用すると、生成AIワークロードの稼働開始までの時間を短縮できます。⁵¹ 詳細については、後の「Dell Validated Designs」のセクションをご覧ください。

Dellはすでに、AIテクノロジー アプリケーションを拡大するためにさまざまな企業とパートナー関係を結んでいます。Hugging Faceと提携して、Dellはサイトでの大規模言語モデル(LLM)のセットアップを容易にしました。このパートナーシップによって、AIに関するHugging Faceの専門知識とDellのサーバーおよびストレージ システムが結合されます。Dell専用のHugging Faceポータルでは、Hugging FaceのオープンソースAIモデルを簡単かつ安全に導入するためのツールが提供されます。その継続的な目標は、それらのモデルをDellのシステムに合わせて改善し、パフォーマンスを向上させ、新しいAIアプリケーションをサポートすることです。⁵²

DellとStarburstは、Starburstの分析にDellのコンピューティングとストレージのテクノロジーを統合したハイパフォーマンスの拡張可能なデータレイク作りを目指していますが、その目的はAIおよびMLツールのためにすべてのデータソースへの単一のアクセス ポイントを提供することです。このパートナーシップを活用すると、データ サイロをなくすことができます。⁵³

当社の調査によると、AIに関するSupermicroのパートナーシップはきわめて限定的です。SiMa.aiとSupermicroが共同開発したSupermicro SYS-E300-13ADは、マルチストリーム ビデオ分析処理用に設計されたコンパクトなエッジMLサーバーです。SiMa.ai MLパイプラインをチップ上に装備するこのサーバーは、複数のビデオ チャネルを効率的に処理し、総所有コストを削減し、信頼性とセキュリティを向上させます。このサーバーの計算に関する設定は、多数のビデオ ストリームの処理と分析のために設計されているため、さまざまなエンタープライズ アプリケーションに適したエッジ インテリジェンスが実現します。⁵⁴

Dell Validated Designs

AIハードウェア ソリューションから当て推量を排除できるように、Dell は、AIおよびその他の多数のワークロードに合わせて最適化された、ラボ検証済みのリファレンス アーキテクチャを提供しています。これらの検証済みの設計には、アーキテクチャの概念、ソリューション全体の概要、パフォーマンスおよびその他のラボ検証が含まれており、ソリューションの設計対象のワークロードに対する能力が実証されています。このワークロードに含まれるものとしては、仮想化環境、MLOps、機械学習、会話型AI、生成AI推論、生成AIモデル チューニング、NVIDIA Fleet Command、OpenShift AIがあります。⁵⁵

検証済み設計の一例である「仮想環境向けAI」は、VMware対応AIと NVIDIA AI EnterpriseをDellインフラストラクチャ上で組み合わせることによって、仮想環境でのAIを最適化するものです。⁵⁶ Validated Design Guideには、ResNetモデルのトレーニング結果を示すパフォーマンス データが含まれているため、設計の有効性が実証されるとともに、顧客がどの種のパフォーマンスを期待できるかが明らかになります。⁵⁷ これらの検証は、単に連携するハードウェアのリストを示すだけでなく、概念を説明し、推奨構成を提示し、考慮事項と期待できるパフォーマンスを詳しく示すという点で価値があります。⁵⁸

Supermicroも、ユースケース ベースでソリューションを提供していますが、Dell Validated Designsのリファレンス アーキテクチャのレベルには達していません。特定のワークロードのための専用ソリューションを推奨する代わりに、Supermicroは自社のサーバーとGPUをAI推論とトレーニング、HPC/AI、可視化と設計などのカテゴリーに分類しています。⁵⁹ 同社のパンフレットとデータシートを見ると、それらのカテゴリーは、該当するタスクに最適であるとして推奨されるいくつかのサーバーとGPUで構成されており、さらにさまざまなユースケースと、関係する主要テクノロジーのリスト、ソフトウェアの候補も含まれています。⁶⁰ Dell Validated Designsとは異なり、ネットワーク アーキテクチャとパフォーマンスまたは検証データは含まれていないようです。Supermicroは、さまざまな「リファレンス デザイン」としていくつかのAIソリューションのための詳細なリファレンス アーキテクチャを提供しています。たとえば、SupermicroとNVIDIAが共同で発行した、水冷を使用する大規模AIトレーニングのソリューション概要です。⁶¹ 顧客の特定のシナリオに適した詳しいリファレンス アーキテクチャが見つかることもありますが、当社の調査の時点で見つかったのは前述の水冷アーキテクチャ、AIワークステーション アーキテクチャ、RedHat OpenShiftアーキテクチャの3件のみでした。⁶²

全体として、Dell Validated Designsは取り上げるAIワークロードの数とガイダンスの詳しさの点でSupermicroのサービスを上回っています。

管理サービスとiDRAC

Principled Technologiesの2023年4月のレポートによると、Integrated Dell Remote Access Controller (iDRAC)は Supermicro Intelligent Platform Management Interface (IPMI)よりも高度な機能を提供しており、特に自動化、セキュリティ、および構成で顕著です。⁶³ 表3は、そのレポートからDellとSupermicroの管理機能の比較をまとめたものであり、iDRACが Supermicro IPMIと比較してどのように導入とファームウェア アップデートが簡単で、多くのセキュリティ機能を提供しているかを示しています。なお、報告された内容の中には当初の発表以降に変化したものもあります。

表3：2023年4月のPrincipled TechnologiesによるDellとSupermicroの管理ツールの比較の要約。報告された内容の中には、発表以降に変化したものもあります。出典：Principled Technologies <https://facts.pt/V5fDf06>

	Dellの管理ツールとの違い	どれだけ優れているか
ファームウェア アップデートの容易性 iDRAC9とSupermicro IPMIの比較 OMEとSupermicro SSMの比較	- iDRAC9ではオンライン アップデートが自動化され、スケジューリングのオプションもある - OMEではカスタム ファームウェア リポジトリの作成が可能であり、BIOS、BMC、その他のサーバー コンポーネントのファームウェアのアップデートに追加のツールやエージェントを必要としない	- iDRACでの自動アップデートの設定に要する時間はわずか74秒 - Supermicro IPMIには自動アップデート機能がないため、管理者は手動でアップデートする必要がある - SSMでサポートされるのはBIOSとBMCのファームウェア アップデートのみであり、その他のコンポーネントのアップデートにはSUMが必要
セキュリティ機能が多い iDRAC9とSupermicro IPMIの比較 OMEとSupermicro SSMの比較	- iDRAC9にはMFAと、システム ダウンタイムのないUSBポート動的無効化機能がある - OMEにはロールベースのアクセス制御(RBAC)とスコープベースのアクセス制御(SBAC)の両方があるため、デバイス管理の対象をデバイス グループのサブセットに限定できる	- Supermicro IPMIにはMFA機能はない - Supermicro IPMIではUSBポートを無効化するにはシステムを再起動してBIOS構成に入る必要がある - Supermicro SSMにはRBACがあるが、それよりも限定的なSBACはない
ライフサイクル管理が簡単 OMEとSupermicro SSMの比較	完全な、エージェントレスのライフサイクル管理がOMEを介して可能であり、管理と監視が容易	SSMでは、ローカル システムの正常性に関する詳細なメトリックを得るにはSuperDoctor5 エージェントが必要であり、追加コンポーネントのアップデートにはSupermicro Update Manager (SUM)が必要
サーバー導入の容易性 iDRAC9とSupermicro IPMIの比較	- 完全なDellサーバー プロファイルのインポートはiDRAC9を使用するとわずか12ステップ - 堅牢なBIOS構成オプションがiDRAC9にあり、52のBIOS機能の設定が可能で、RAID NICやiDRACなどのコンポーネント構成もサポートされる	- Supermicro IPMIではサーバー プロファイル全体ではなくIPMI構成の保存と復元のみが可能 - iDRAC9には52のBIOS機能があるが、IPMIにはBIOS構成のオプションはない
レポート作成と分析に関するオプションが多い iDRAC9とSupermicro IPMIの比較 OMEとSupermicro SSMの比較	- iDRAC9にはテレメトリ ストリーミングがあり、サーバー データをSplunkなどの分析ツールに簡単に送信できる - OMEによってテレメトリ データがCloudIQに直接送信されるため、監視が容易	- IPMIにあるのは集計と最終的な分析のために管理者がメッセージを送信するのに使用できるSYSLOG機能のみ - SSMにはDell CloudIQと同等のクラウドベースの管理ソリューションはない
サステナビリティに関する機能が多い OMEとSupermicro SSMの比較	OME Power Managerに監視のための多数のメトリックがある（二酸化炭素排出量のデータも含まれる）	SSMは堅牢な使用率メトリックが少なく、二酸化炭素排出量を追跡する方法はない
監視方法が多い OMEとSupermicro SSMの比較	- Dellサーバーの管理がどこからでもOpenManage Mobile アプリケーションでできる - サードパーティー製デバイスの監視を、サーバーIPと認証情報を使用してOMEで行うことができ、サードパーティーSNMP MIBのインポートもサポートされる	- SSMにはモバイル アプリケーションはない - SSMではサードパーティー製デバイスの監視をサーバーIPを使用して行うことは許可されない



結論

ビジネスにおけるAIソリューションの設計、実装、管理、維持に関しては、考慮すべき多くの要素があります。賢明に投資してAIソリューションを最大限に活用するためには、ハードウェア サプライヤー以上の存在になれるベンダーを探すことをお勧めします。本レポートの調査結果が示すように、Dellはパートナーとして、お客様の導入過程全体をサポートするサービスを提供しているため、AIを始めるにあたってはDellへの投資をぜひご検討ください。

1. Vipera「NVIDIA's H100 and A100 GPU Cards: Exploring the Intricacies of SXM and PCI-E Connections」(2024年1月5日にアクセス) <https://www.viperatech.com/unraveling-the-mysteries-sxm-vs-pci-e-connections-in-nvidias-high-end-h100-and-a100-gpus/>
2. MLCommons「MLPerf Inference: Datacenter Benchmark Suite Results」(2024年1月5日にアクセス) <https://mlcommons.org/en/inference-datacenter-31/>
3. MLCommons「MLPerf Inference: Datacenter Benchmark Suite Results」
4. Dell「PowerEdge XE Servers」(2024年1月5日にアクセス) <https://www.dell.com/en-us/dt/servers/specialty-servers/poweredge-xe-servers.htm>
5. Supermicro「Next Leap of AI Infrastructure is Here」(2024年1月5日にアクセス) <https://www.supermicro.com/en/accelerators/nvidia>
6. Dell「PowerEdge XE9680」(2024年1月5日にアクセス) <https://www.delltechnologies.com/asset/en-us/products/servers/technical-support/poweredge-xe9680-spec-sheet.pdf>
7. 検証済みMLPerf®スコアv3.1推論、Closed。 <https://mlcommons.org/benchmarks/inference-datacenter/> (2023年12月5日、エントリー3.1-0069から抽出)。MLPerfの名前とロゴは、米国およびその他の国におけるMLCommons Associationの登録商標および未登録商標です。All rights reserved. (不許複製・禁無断転載) 無断での使用は固く禁じられています。詳細については、www.mlcommons.orgを参照してください
8. 検証済みMLPerf®スコアv3.1推論、Closed。 <https://mlcommons.org/benchmarks/inference-datacenter/> (2023年12月5日、エントリー3.1-0135から抽出)。MLPerfの名前とロゴは、米国およびその他の国におけるMLCommons Associationの登録商標および未登録商標です。All rights reserved. (不許複製・禁無断転載) 無断での使用は固く禁じられています。詳細については、www.mlcommons.orgを参照してください
9. 検証済みMLPerf®スコアv3.1推論、Closed。 <https://mlcommons.org/benchmarks/inference-datacenter/> (2023年12月5日、エントリー3.1-0069から抽出)。MLPerfの名前とロゴは、米国およびその他の国におけるMLCommons Associationの登録商標および未登録商標です。All rights reserved. (不許複製・禁無断転載) 無断での使用は固く禁じられています。詳細については、www.mlcommons.orgを参照してください
10. 検証済みMLPerf®スコアv3.1推論、Closed。 <https://mlcommons.org/benchmarks/inference-datacenter/> (2023年12月5日、エントリー3.1-0132から抽出)。MLPerfの名前とロゴは、米国およびその他の国におけるMLCommons Associationの登録商標および未登録商標です。All rights reserved. (不許複製・禁無断転載) 無断での使用は固く禁じられています。詳細については、www.mlcommons.orgを参照してください

11. Dell「PowerEdge XE9640 Rack Server」(2024年1月5日にアクセス)
<https://www.dell.com/en-us/shop/ipovw/poweredge-xe9640>
12. Accelsius「Enabling the AI Revolution with Liquid Cooling」(2024年1月5日にアクセス)
<https://www.accelsius.com/blog/enabling-the-ai-revolution-with-liquid-cooling>
13. Dell「PowerEdge XE9640 Technical Guide」(2024年1月5日にアクセス) <https://www.delltechnologies.com/asset/en-us/products/servers/technical-support/poweredge-xe9640-technical-guide.pdf>
14. Supermicro「GPU Server Systems」(2024年1月5日にアクセス)
https://www.supermicro.com/en/products/gpu?pro=pl_grp_type%3D1
15. 検証済みMLPerf[®]スコアv3.1推論、Closed。 <https://mlcommons.org/benchmarks/inference-datacenter/> (2023年12月5日、エントリー3.1-0067から抽出)。MLPerfの名前とロゴは、米国およびその他の国におけるMLCommons Associationの登録商標および未登録商標です。All rights reserved. (不許複製・禁無断転載) 無断での使用は固く禁じられています。詳細については、www.mlcommons.orgを参照してください
16. 検証済みMLPerf[®]スコアv3.1推論、Closed。 <https://mlcommons.org/benchmarks/inference-datacenter/> (2023年12月5日、エントリー3.1-0133から抽出)。MLPerfの名前とロゴは、米国およびその他の国におけるMLCommons Associationの登録商標および未登録商標です。All rights reserved. (不許複製・禁無断転載) 無断での使用は固く禁じられています。詳細については、www.mlcommons.orgを参照してください
17. Dell「PowerEdge XE8640」(2024年1月5日にアクセス)
<https://www.delltechnologies.com/asset/en-us/products/servers/technical-support/poweredge-xe8640-spec-sheet.pdf>
18. Dell「PowerEdge XE8640 Rack Server」(2024年1月5日にアクセス)
<https://www.dell.com/en-us/shop/ipovw/poweredge-xe8640>
19. 検証済みMLPerf[®]スコアv3.1推論、Closed。 <https://mlcommons.org/benchmarks/inference-datacenter/> (2023年12月5日、エントリー3.1-0066から抽出)。MLPerfの名前とロゴは、米国およびその他の国におけるMLCommons Associationの登録商標および未登録商標です。All rights reserved. (不許複製・禁無断転載) 無断での使用は固く禁じられています。詳細については、www.mlcommons.orgを参照してください
20. 検証済みMLPerf[®]スコアv3.1推論、Closed。 <https://mlcommons.org/benchmarks/inference-datacenter/> (2023年12月5日、エントリー3.1-0133から抽出)。MLPerfの名前とロゴは、米国およびその他の国におけるMLCommons Associationの登録商標および未登録商標です。All rights reserved. (不許複製・禁無断転載) 無断での使用は固く禁じられています。詳細については、www.mlcommons.orgを参照してください
21. Dell「Artificial Intelligence (AI) technologies powered by Dell Precision workstations」(2024年1月5日にアクセス)
<https://www.dell.com/en-us/dt/ai-technologies/index.htm?hve=explore+dell+precision+for+ai#tab0=0>
22. Supermicro「Super Workstations」(2024年1月5日にアクセス) <https://www.supermicro.com/en/products/superworkstation>
23. Supermicro「Rackmount Workstations」(2024年1月5日にアクセス)
<https://www.supermicro.com/en/products/rackmount-workstations>
24. Stephen Pritchard「Storage requirements for AI, ML and analytics in 2022」(2024年1月5日にアクセス)
<https://www.computerweekly.com/feature/Storage-requirements-for-AI-ML-and-analytics-in-2022>
25. Dell「PowerScale AI-Ready Data Platform」(2024年1月5日にアクセス)
<https://www.dell.com/en-us/shop/powerscale-family/sf/powerscale>
26. Dell「Dell PowerScale All-Flash」(2024年01月05日にアクセス) <https://www.delltechnologies.com/asset/en-us/products/storage/technical-support/h15963-ss-powerscale-all-flash-nodes.pdf>
27. Dell「Dell PowerScale OneFS Software Features」(2024年01月05日にアクセス) <https://www.delltechnologies.com/asset/en-us/products/storage/technical-support/h18275-onefs-software-features-data-sheet.pdf>
28. Dell「Dell ECS Enterprise Object Storage」(2024年1月5日にアクセス) <https://www.dell.com/en-us/dt/storage/ecs/>
29. Supermicro「Accelerating AI Data Pipelines」(2024年1月5日にアクセス) <https://www.supermicro.com/en/products/storage>
30. Supermicro「Supermicro Software-Defined Storage and Memory Solutions」(2024年1月5日にアクセス)
<https://www.supermicro.com/en/solutions/software-defined-storage>
31. Supermicro「Supermicro WEKA Distributed Storage Solution」(2024年1月5日にアクセス)
<https://www.supermicro.com/en/solutions/wekaio>
32. Dell「Dell ECS Enterprise Object Storage」(2024年1月5日にアクセス) <https://www.dell.com/en-us/dt/storage/ecs/>
33. Dell「Dell ECS Enterprise Object Storage」
34. Supermicro「Supermicro OSNEXUS Software-Defined Storage Solution」(2024年1月5日にアクセス)
<https://www.supermicro.com/en/solutions/osnexus>

-
35. ASBIS「Supermicro Solution for Scalify RING」(2024年1月5日にアクセス)
<https://news.asbis.com/news/suppliers/supermicro-renewed-the-line-of-scalify-ring-solution/>
 36. Supermicro「Supermicro solution for Quantum ActiveScale」(2024年1月5日にアクセス)
<https://www.supermicro.com/en/solutions/activescale>
 37. Dell「Scalable and elastic Storage as-a-Service」(2024年1月5日にアクセス)
<https://www.dell.com/en-us/dt/apex/storage/data-storage-services/>
 38. Dell「Flip the Switch to Open Networking with PowerSwitch」(2024年1月5日にアクセス)
<https://www.dell.com/en-us/dt/networking/>
 39. Dell「Dell PowerSwitch Data Center Switches」(2024年1月5日にアクセス)
<https://www.dell.com/en-us/dt/networking/data-center-switches/>
 40. Supermicro「SSE-T7132S - 400Gb Ethernet Switch」(2024年1月5日にアクセス)
<https://www.supermicro.com/en/products/accessories/Networking/SSE-T7132SR.php>
 41. Dell「Dell EMC Networking SmartFabric Services Deployment with VxRail 4.7—Fabric Design Center」(2024年1月5日にアクセス) <https://infohub.delltechnologies.com/l/dell-emc-networking-smartfabric-services-deployment-with-vxrail-4-7-1/fabric-design-center-26/>
 42. Dell「Dell SmartFabric Services」(2024年1月5日にアクセス) <https://www.dell.com/en-us/dt/networking/smartfabric/>
 43. Penny Madsen「Scaling Skills for AI: Lessons from Early Adopters」(2024年1月5日にアクセス) <https://www.delltechnologies.com/asset/en-us/products/servers/industry-market/idc-brief-importance-of-skills-for-ai-dell.pdf>
 44. Dell「Design Guide—Generative AI in the Enterprise – Model Customization—Overview」(2024年1月5日にアクセス)
<https://infohub.delltechnologies.com/l/design-guide-generative-ai-in-the-enterprise-model-customization/overview-5381/>
 45. Dell「Design Guide—Generative AI in the Enterprise – Model Customization—Advisory Services for Generative AI」(2024年1月5日にアクセス) <https://infohub.delltechnologies.com/l/design-guide-generative-ai-in-the-enterprise-model-customization/advisory-services-for-generative-ai-1/>
 46. Dell「Design Guide—Generative AI in the Enterprise – Model Customization—Adoption Services for Generative AI」(2024年1月5日にアクセス) <https://infohub.delltechnologies.com/l/design-guide-generative-ai-in-the-enterprise-model-customization/adoption-services-for-generative-ai-1/>
 47. Dell「Design Guide—Generative AI in the Enterprise – Model Customization—Adoption Services for Generative AI」
 48. Dell「Artificial Intelligence (AI) Ready Solution Services」(2024年1月5日にアクセス)
<https://www.dell.com/en-us/dt/services/solutions/artificial-intelligence-services.htm>
 49. Dell「Comprehensive AI Training Modules Tailored for You」(2024年1月5日にアクセス)
<https://education.dell.com/content/emc/en-us/home/training/aiml.html>
 50. Supermicro「Services and Support」(2024年1月5日にアクセス) <https://www.supermicro.com/en/support>
 51. Travis Vigil「Dell and NVIDIA: Bringing Generative AI to the Enterprise」(2024年1月5日にアクセス)
<https://www.dell.com/en-us/blog/dell-and-nvidia-bringing-generative-ai-to-the-enterprise/>
 52. Dell「Dell Technologies and Hugging Face to Simplify Generative AI with On-Premises IT」(2024年1月5日にアクセス) <https://www.dell.com/en-us/dt/corporate/newsroom/announcements/detailpage.press-releases~usa~2023~11~20231114-dell-technologies-and-hugging-face-to-simplify-generative-ai-with-on-premises-it.htm>
 53. Richard DeMare「Starburst and Dell expand partnership to accelerate AI efforts with more intelligent data collection」(2024年1月5日にアクセス) <https://www.starburst.io/blog/starburst-and-dell-expand-partnership-to-accelerate-ai-efforts-with-more-intelligent-data-collection/>
 54. Business Wire「SiMa.ai and Supermicro Announce Partnership to Accelerate Power Efficient ML at the Edge」(2024年1月5日にアクセス) <https://www.businesswire.com/news/home/20231129609794/en/SiMa.ai-and-Supermicro-Announce-Partnership-to-Accelerate-Power-Efficient-ML-at-the-Edge>
 55. Dell「Dell AI Solutions」(2024年1月5日にアクセス)
<https://www.dell.com/en-us/dt/solutions/artificial-intelligence/index.htm#accordion0&tab0=0>
 56. Dell「Unlock the power of AI in virtualized environments」(2024年1月5日にアクセス)
<https://www.delltechnologies.com/asset/en-us/products/ready-solutions/briefs-summaries/ai-vxrail-powerscale-brief.pdf>
 57. Dell「Design Guide—Virtualizing GPUs for AI with VMware and NVIDIA Based on Dell Infrastructure—Performance Results」(2024年1月5日にアクセス) <https://infohub.delltechnologies.com/l/design-guide-virtualizing-gpus-for-ai-with-vmware-and-nvidia-based-on-dell-infrastructure-1/performance-results-15/>

-
58. Dell「Design Guide—Virtualizing GPUs for AI with VMware and NVIDIA Based on Dell Infrastructure—Design Considerations」(2024年1月5日にアクセス) <https://infohub.delltechnologies.com//design-guide-virtualizing-gpus-for-ai-with-vmware-and-nvidia-based-on-dell-infrastructure-1/design-considerations-105/>
59. Supermicro「Accelerate Every Workload」(2024年1月5日にアクセス) <https://www.supermicro.com/en/solutions/ai-deep-learning>
60. Supermicro「Supermicro Enterprise AI Inference & Training」(2024年1月5日にアクセス) https://www.supermicro.com/datasheet/Datasheet_AI-Workloads_Enterprise_AI_Inferencing_and_Training.pdf
61. Supermicro「SUPERMICRO RACK SCALE SOLUTIONS: LARGE SCALE AI TRAINING WITH LIQUID COOLING」(2024年1月5日にアクセス) https://www.supermicro.com/solutions/Solution-Brief_Rack_Scale_AI.pdf
62. Supermicro「Accelerate Every Workload」(2024年1月5日にアクセス) <https://www.supermicro.com/en/solutions/ai-deep-learning>
63. Principled Technologies「Dell management tools made server deployment and updates easier, offered more comprehensive security, and provided more robust infrastructure analytics」(2024年1月5日にアクセス) <https://www.principledtechnologies.com/Dell/Management-tools-vs-Supermicro-0423.pdf>

▶ 本レポートのオリジナルの英語版はこちらでご確認いただけます : <https://facts.pt/q9p46K9>

このプロジェクトは、[デル・テクノロジーズ]の委託を受けて作成されています。



Facts matter.®

Principled Technologiesは、Principled Technologies, Inc.の登録商標です。他のすべての製品名は各社の商標です。

保証の免責事項、責任の制限：

Principled Technologies, Inc.はそのテストの精度と妥当性を確保するために適切な努力を行っていますが、テストの結果と分析、それらの精度、完全性、または品質に関して、特定の目的に対する適合性の黙示保証を含め、明示または黙示にかかわらず、いかなる保証も放棄します。すべての個人または事業体は自己の責任においてテストの結果に依存し、Principled Technologies, Inc.およびその従業員、その請負業者が、テスト手順や結果における疑わしいエラーや欠陥による損失や損害についてのいかなる主張に対しても、何ら責任を負わないことを認めるものとします。

Principled Technologies, Inc.は、そのテストに関連する間接的、特別的、付随的、結果的な損害に対して、当該損害の可能性について知らされていた場合でも、一切責任を負わないものとします。いかなる場合もPrincipled Technologies, Inc.は、直接的損害を含め、Principled Technologies, Inc.のテストに関連して支払われた金額を超える責任を負わないものとします。お客様の唯一の救済手段は、ここに示すとおりです。