



Dell AI 포트폴리오로 AI 워크로드의 당면 과제 해결

Dell AI 포트폴리오와 HPE의 유사 오퍼링 비교

AI(Artificial Intelligence)가 비즈니스 환경을 혁신했으며 모든 규모의 산업과 조직이 데이터에서 더 심층적인 통찰력을 얻고, 비즈니스 프로세스를 자동화하며, 맞춤형 고객 및 사용자 경험을 제공하고, 업계에서 경쟁력을 강화할 수 있도록 지원했다는 것은 의심의 여지가 없다. AI의 능력을 효과적으로 활용하려면 조직은 전체 AI 수명주기를 아우르는 포괄적인 통합 솔루션 포트폴리오를 제공할 수 있는 인프라스트럭처 공급업체가 필요하다.

Dell Technologies 및 HPE(Hewlett Packard Enterprise)와 같은 인프라스트럭처 공급업체는 고객이 증가하는 AI 수요에 대응하고 AI의 본질적인 복잡성을 극복할 수 있도록 지원한다. AI 지원 포트폴리오를 제시하는 이러한 공급업체는 고성능 온프레미스 및 클라우드 인프라스트럭처 솔루션과 전략적 파트너십, 다양한 지원 및 컨설팅 서비스를 통합하는 다양한 수준의 AI 솔루션을 제공한다.

이 보고서에서는 고객이 AI 요구 사항에 맞게 Dell Technologies를 선택할 경우 얻을 수 있는 구체적인 아키텍처, 성능 및 지원 이점을 강조하기 위해 Dell 및 HPE AI 포트폴리오*에 대해 공개된 정보를 검토한다. Dell이 AI 배포를 지원하기 위해 구축한 서버의 세부 정보와 ML Commons®의 기준 업계 벤치마크 테스트 결과를 비교한다. 또한 AI 여정의 모든 단계에서 고객을 지원하는 추가 소프트웨어 및 서비스 오퍼링도 살펴본다.

*참고: PT는 2023년 12월 5일 또는 그 이전에 모든 조사를 완료했으므로, 해당 날짜 이후 Dell 또는 HPE 릴리스의 오퍼링 또는 변경 사항은 이 백서에 반영되지 않는다.

AI 도입의 당면 과제

AI 전략을 채택하면 데이터 센터와 데이터 센터에 근무하는 IT 인력은 다음과 같은 수많은 새로운 과제에 직면하게 된다.

- 사내 AI 교육 또는 외부 채용을 통해 현재 직원의 기존 기술 격차 해소
- 비즈니스 데이터의 품질, 수량, 위치 및 현재 상태를 포함하여 AI의 데이터 준비 요구 사항 이해
- 특정 비즈니스 AI 목표를 진단하여 어떤 AI 모델 및 구현이 유익할지 더 현명하게 판단
- 계획된 AI 시스템의 컴퓨팅, 네트워킹 및 스토리지 요구 사항을 진단하고 획득 계획 결정

이는 기업이 데이터 센터에 AI를 구현함으로써 이점을 얻고자 할 때 직면하게 되는 거대한 장애물 중 몇 가지 예일 뿐이다.

Dell AI 포트폴리오는 고객이 구현 로드맵을 구축하고 AI 모델을 위한 데이터를 준비하는 데 도움이 되는 전문 서비스와 컨설팅 서비스를 통해 고객이 이러한 과제를 해결하도록 지원한다.¹ 또한 이 포트폴리오는 ML(Machine Learning) 개념 및 기타 교육 주제를 다루는 교육 과정을 포함하고 있으며 성공적인 구현을 위해 Validated Designs for AI를 제공한다.² 더 나아가 Dell Technologies는 타사와 협력하여 고객에게 추가 AI 툴을 제공하는데, 여기에는 오픈 소스 AI 모델 배포를 위한 전용 컨테이너와 스크립트가 포함된 Hugging Face 커뮤니티 내의 맞춤형 Dell 포털³과 Meta Llama 2 LLM(Large Language Model)의 간편한 배포가 포함된다.⁴ Dell Technologies는 모바일 워크스테이션부터 최대 8개의 하이엔드 NVIDIA GPU를 지원하는 서버에 이르기까지 광범위한 컴퓨팅 및 PC 오퍼링과 함께 고성능 파일 및 오브젝트 스토리지 어레이 포트폴리오를 통해 AI에 필요한 비정형 데이터 스토리지도 제공한다. Dell PowerScale, ObjectScale, ECS 및 온보드 스토리지를 비롯한 이러한 스토리지 오퍼링은 AI 워크로드에 자주 사용되는 비정형 데이터를 처리할 수 있다.⁵ 또한 Dell Technologies는 Snowflake와 협력하여 Dell 고객에게 하이브리드 클라우드 스토리지 솔루션을 제공하였다.⁶ Dell Technologies의 분석에 따르면, 2023년 8월 현재 Dell Technologies는 AI 서버와 스토리지를 넘어 구현 여정 전반에 걸쳐 리소스를 제공함으로써 "광범위한 Generative AI 포트폴리오"를 제공한다.⁷

AI 성능 및 가속화된 컴퓨팅 옵션: Dell과 HPE 비교

AI 워크로드는 워크로드의 크기나 유형에 따라 CPU나 GPU 또는 두 가지 모두를 컴퓨팅 리소스로 사용할 수 있다. 일부 CPU는 최신 인텔 제온 스케일러블 프로세서의 인텔 AMX(Advanced Matrix Extensions)와 같은 AI 전용 가속기를 제공한다.⁸ 많은 경우 GPU는 규모가 더 크거나 더 까다로운 워크로드에 적합하지만 GPU 폼 팩터가 성능 수준에 영향을 미칠 수 있다. 예를 들어 일부 NVIDIA A100 및 H100 모델 GPU는 범용 PCIe 또는 독점 SXM 폼 팩터로 제공되며, 그중 후자는 고성능 NVIDIA SXM 아키텍처를 사용한다.⁹ 큰 메모리 용량과 발열 억제 아키텍처 및 전력 효율성과 같은 서버 설계 기능도 성능에 영향을 미친다. 대부분의 데이터 센터에서는 여전히 공기 냉각을 사용한다. 즉, HPC(High-Performing Compute) 워크로드에는 가능한 한 효과적으로 공기를 이용해 냉각하도록 제작된 서버가 필요하다. 아래에서는 공개된 MLCommons® MLPerf® 점수와 함께 구성 요소, 냉각 옵션 등의 측면에서 PowerEdge 서버 오퍼링을 중점적으로 살펴본다.

AI 모델 벤치마크 성능: MLPerf 결과 비교

MLPerf®는 학습과 추론 모두에 대한 AI 성능을 테스트하는 벤치마크 제품군이다. 조직이 공식 MLPerf® 결과를 게시하려면 벤치마크 개발자인 MLCommons®가 설정한 특정 조건에 결과가 부합해야 한다.¹⁰ 이러한 규정 준수 지침은 성능 비교를 더 용이하게 하는 표준을 제공한다. 추론 테스트의 경우 MLPerf®는 데이터 센터, 엣지, 모바일 및 소규모 데이터 세트를 사용하고 테스트 중에 소비된 전력의 와트수와 AI 점수를 보고한다. 추론 벤치마크 제품군에는 여러 일반적인 AI, ML 및 DL 모델에 대한 테스트가 포함되어 있다(표 1 참조).

1: MLPerf® 테스트에 포함된 AI, ML 및 DL 모델과 각 모델의 일반적인 활용 사례. 출처: Principled Technologies.

일반적인 AI 모델	일반적인 활용 사례
ResNet	컴퓨터가 의료 영상, 소셜 미디어 콘텐츠 관리, 안면 인식과 같은 활용 사례를 위한 다양한 이미지를 학습, 기억 및 식별하는 데 도움이 되는 이미지 분류 모델
RetinaNet	ResNet에 비해 추가적인 복잡성을 처리할 수 있는 오브젝트 감지 유형. 컴퓨터가 이미지 또는 비디오 프레임 내에서 오브젝트를 식별하고 찾는 데 도움이 되며 중요도별로 분류할 수 있다. 자율 주행, 차량 자동 지원 기술, 감시, 안면 인식과 같은 용도로 사용된다.
3D-UNet	의료 영상 세분화에만 해당
RNN-T	자동 언어 번역과 같은 활용 사례를 위한 음성 인식
BERT	텍스트 요약, 언어 번역, 작업 자동 완성 등의 활용 사례를 위한 자연어 처리
DLRM-v2-99.9	타겟 광고, 맞춤형 제품 추천 등의 활용 사례를 위한 권장 모델
GPTJ-99 및 99.9	챗봇 및 채팅 기반 AI 톨과 같은 활용 사례를 위한 텍스트 생성에 탁월한 자연어 처리용 LLM

MLPerf

MLPerf® 결과에는 AI 모델 자체 외에도 여러 매개변수가 포함되어 있으므로 단일 차트 또는 표에서 많은 데이터를 구문 분석해야 할 수 있다. 다음은 이러한 매개변수에 대한 빠른 참조이다.

- 99.0 및 99.9: 이러한 수치는 모델이 학습된 정확도를 나타낸다. 출력이 더 정확해야 할수록 모델이 더 복잡해지고 데이터 처리 시간이 길어질 수 있다.
- Offline samples/sec: 벤치마크가 테스트 시작 시 모든 쿼리를 전송하여 시스템에 이미 있는 데이터를 시뮬레이션하는 모드
- Server queries/sec: 벤치마크가 테스트 기간 내내 쿼리를 전송하여 라이브 데이터 스트림 분석을 시뮬레이션하는 모드

MLCommons® 및 MLPerf® 결과에 대한 자세한 내용은 <https://mlcommons.org/benchmarks/inference-datacenter/>에서 참조할 수 있다.

이 보고서의 결과는 2023년 11월부터 MLCommons® 웹사이트에 게시된 MLPerf® v3.1 Inference Datacenter 결과에서 발췌한 것이다.¹¹ 이러한 결과에는 기술 제조업체와 클라우드 서비스 공급업체의 제출물이 포함되며 다양한 구성이 포괄되어 있다. HPE에서 공개적으로 제출한 자료에 비해 Dell 서버는 특정 AI 모델에서 더 나은 결과를 보였다. (참고: 서버마다 GPU 구성이 다르면 직접 비교가 어려울 수 있다.) 자세한 내용은 표 2에서 확인할 수 있다.

2: 2023년 11월 29일 현재 게시된 MLCommons® MLPerf® 3.1 결과에 포함된 Dell 및 HPE 서버.
출처: Principled Technologies.

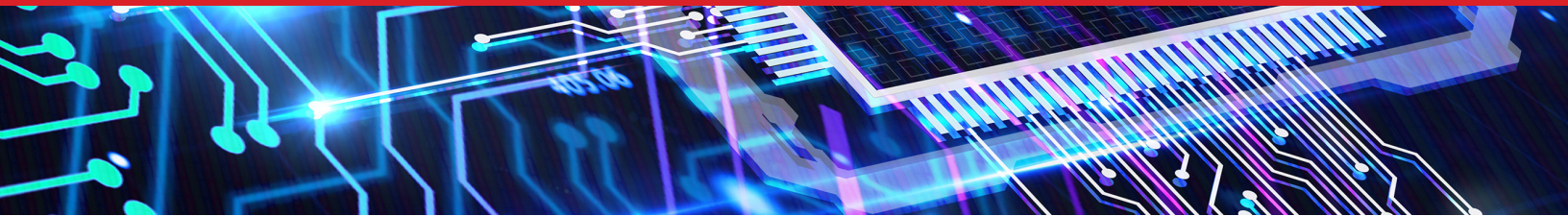
제출자	서버 모델	GPU 수 및 모델	설명
Dell ¹²	PowerEdge XE9680	8x NVIDIA H100 SXM	대규모 언어 모델과 같은 대규모 워크로드에서 AI 학습 및 추론용
	PowerEdge XE9640	4x NVIDIA H100 SXM	고밀도 및 액체 냉각식 데이터 센터에서 대규모 AI 모델 학습용
	PowerEdge XE8640	4x NVIDIA H100 SXM	공기 냉각식 데이터 센터를 위한 4U 폼 팩터에서 기존 AI 학습, HPC 및 데이터 분석 애플리케이션 구동용
	PowerEdge R760xa	4x NVIDIA H100 PCIe	최고 성능의 GPU가 필요하지 않은 AI-ML/DL 학습 및 추론을 비롯한 광범위한 고성능 컴퓨팅 워크로드용
HPE ^{13,14}	ProLiant XL675d Gen10 Plus	8x NVIDIA A100 SXM	고성능 컴퓨팅 및 AI용
	ProLiant DL380a Gen11	4x NVIDIA H100 PCIe	보통 수준의 AI 워크로드를 위한 2U 서버

Dell과 HPE 서버 간의 직접 비교

하드웨어보다 포괄적인 AI 전략이 더 중요하지만, 가장 강력한 하드웨어 성능을 보장하는 것은 AI 워크로드의 성공에 중요한 요소 중 하나이다. 최신 GPU 및 기타 기술을 사용할 수 있게 됨에 따라 AI 워크로드 기능도 발전한다. MLPerf® v3.1 결과가 처음 공개되었을 당시, 사용 가능한 최고의 NVIDIA GPU는 H100 Tensor Core였으며, Dell Technologies는 PCIe 및 SXM5 폼 팩터 모두에서 여러 서버에 대한 MLPerf® 결과를 공개했다.¹⁵ 공개된 HPE 결과에는 H100 제출물이 하나뿐이었고 PCIe 폼 팩터만 포함되었다. 조사에 따르면, 사용 가능한 GPU 지원 HPE 서버 중 최상의 NVIDIA GPU 성능을 위해 SXM5 H100 폼 팩터를 지원하는 서버는 거의 없으며 HPE ProLiant 서버 중에서는 이를 지원하는 서버가 하나도 없다.¹⁶ 아래에서 볼 수 있듯이, GPU 성능이 높을수록 일반적으로 AI 워크로드 성능이 향상된다.

8개 GPU에 대한 MLPerf 결과

Dell PowerEdge XE9680은 AI 가속을 위한 최대 8개의 NVIDIA H100 SXM5 GPU와 최대 2개의 4세대 인텔® 제온® 스케일러블 프로세서에 대한 지원을 제공한다. PowerEdge XE 제품군에는 AMD의 SXM4 또는 SXM5 NVIDIA GPU 또는 OAM(Open Compute Project Accelerator Module) GPU 어셈블리를 지원하는 모듈형 아키텍처가 있어 표준 PCIe GPU에 비해 성능을 강화할 수 있다.¹⁷ 6U의 랙 공간만 차지하는 PowerEdge XE9680은 컴팩트한 8웨이 NVIDIA H100 SXM5 서버이다. 최신 HPE ProLiant Gen11 서버는 현재 H100 SXM 폼 팩터를 지원하지 않지만¹⁸ 일부 HPE Cray Supercomputing 서버는 지원한다.¹⁹ HPE는 Cray 서버에 대한 MLPerf® 결과를 제출하지 않았고 AI 포트폴리오 페이지에서 ProLiant 서버만 강조하므로 이 백서에서는 ProLiant 서버에 중점을 두겠다. (그림 1 참조.)



Featured AI products and services

PRODUCT

HPE Ezmeral Unified Analytics Software

Unlock data and insights faster by helping you develop and deploy data and analytic workloads. Provides fully managed, secure, enterprise-grade versions of the most popular open-source frameworks with a consistent SaaS experience.

[Explore more →](#)

PRODUCT

HPE Machine Learning Development Environment

Uncover hidden insights from your data by helping engineers and data scientists collaborate, build more accurate ML models and train them faster.

[Explore more →](#)

PRODUCT

HPE Machine Learning Data Management Software

Uncover hidden insights with a data pipelining and versioning solution that automates data pipelines and accelerates time to ML model production by processing petabyte-scale workloads.

[Explore more →](#)

PRODUCT

HPE ProLiant Servers

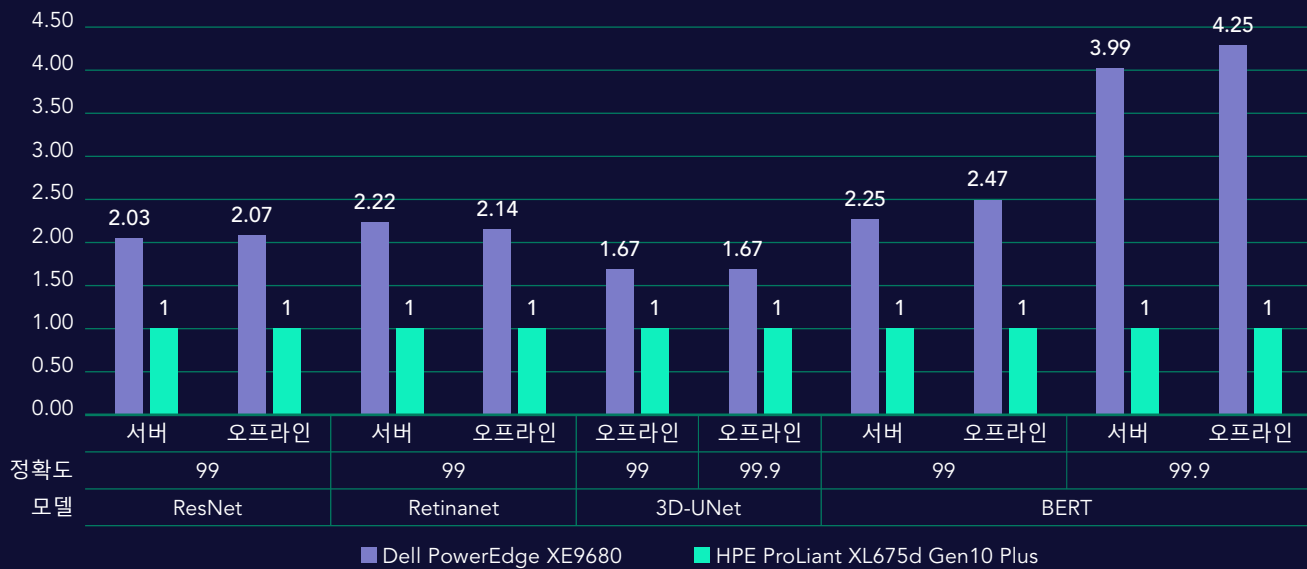
Speed time to value with systems that are optimized for computer vision inference, generative visual AI, and end-to-end natural language processing.

[Explore more →](#)

1: 2023년 12월 5일 현재 <https://www.hpe.com/us/en/solutions/ai-artificial-intelligence.html>에서 HPE ProLiant 서버를 강조하고 있는 주요 AI 제품 및 서비스의 스크린샷.

2023년 11월에 처음 공개된 8개 GPU 서버에 대한 MLPerf® v3.1 결과에서 NVIDIA SXM5 H100 GPU를 탑재한 Dell PowerEdge XE9680은 NVIDIA SXM4 A100 GPU를 탑재한 HPE ProLiant XL675d Gen10 Plus보다 최대 4.25배 더 우수한 성능을 보였다(그림 2 참조).

정규화된 MLPerf® 결과: H100 SXM5가 탑재된 Dell PowerEdge XE9680과 A100 SXM4가 탑재된 HPE ProLiant XL675d Gen10 Plus 비교 (클수록 더 우수)



2: 2023년 11월 29일 현재 Dell PowerEdge XE9680 및 HPE ProLiant XL675d Gen10 Plus에 대해 공개된 MLPerf® 결과. Dell 시스템은 NVIDIA H100 GPU를 사용하지만 HPE 시스템의 GPU는 한 세대 이전 버전이다. 출처: MLCommons®의 데이터를 사용하는 Principled Technologies.^{20,21}

쉽게 비교할 수 있도록 그림 2~5의 테스트 결과를 정규화했다. 즉, 각 HPE ProLiant DL380a Gen 11 결과에 1의 값을 할당하고 이와 관련된 Dell PowerEdge R760xa 결과를 표시하는 것이다. 이러한 결과에서 알 수 있듯이, GPU 모델 간에 단 한 세대 차이만 있어도 다양한 AI 워크로드에서 예상할 수 있는 성능에 상당한 차이가 생길 수 있다.

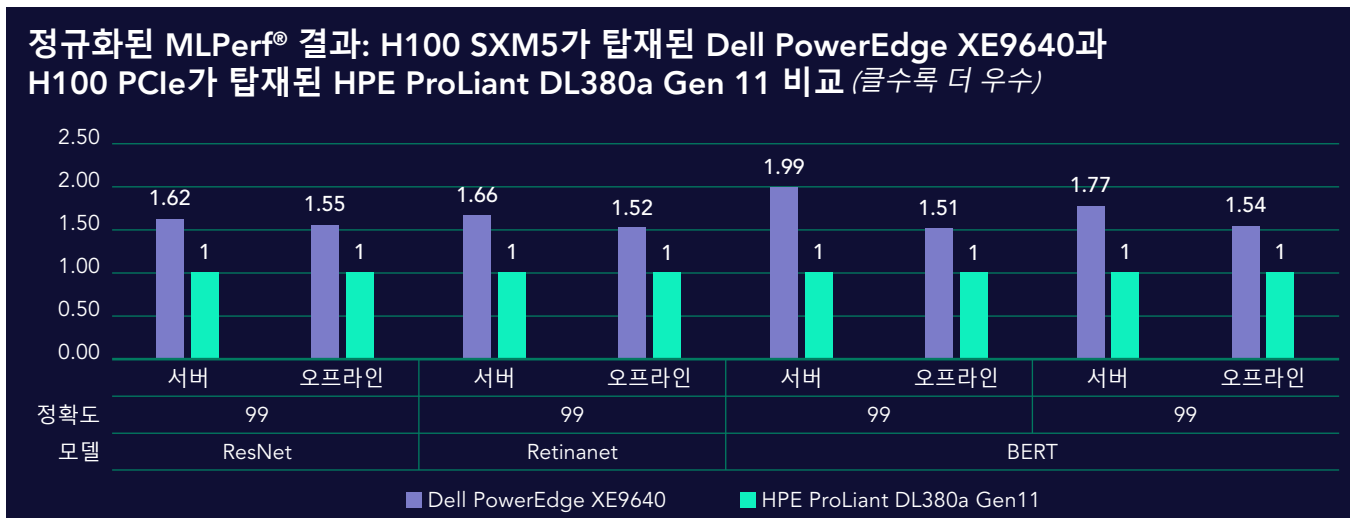
4개 GPU에 대한 MLPerf 결과

데이터 센터 전력 또는 공간 절약이 주요 관심사라면 2U Dell PowerEdge XE9640이 해답을 제공할 수 있다. 최대 4개의 NVIDIA H100 SXM GPU를 탑재한 PowerEdgeXE9640은 XE9680의 절반에 달하는 GPU 컴퓨팅 성능을 2/3 더 작은 공간에서 제공한다.²² 고밀도로 탑재된 Dell PowerEdge XE9640은 Dell Smart Cooling 기술을 통합하여 CPU 및 GPU를 위한 직접 액체 냉각을 포함하여 다양한 방열 기술을 제공한다.²³

PowerEdge XE9640의 2U 새시는 PCIe 카드 및 메모리와 같은 기타 중요한 구성 요소를 냉각할 수 있도록 더 큰 팬과 방열판을 비롯한 향상된 공기 흐름 메커니즘을 수용한다.²⁴ PowerEdge XE9640은 현재 Dell 또는 HPE에서 제공하는 오퍼링 중 2U에서 4웨이 HGX H100 GPU를 탑재한 유일한 오퍼링이다. HPE AI 포트폴리오는 1U 및 2U Gen11 ProLiant 서버를 제공하지만 PCIe 폼 팩터 GPU로 제한된다.²⁵

Dell PowerEdge XE9640 서버는 PCIe 카드와 OAM(OpenCompute Accelerator Module)을 포함하는 저전력, 고밀도 GPU 옵션을 제공하는 인텔 Max GPU 시리즈 1550 OAM GPU도 지원한다.²⁶ 2023년 12월 5일 현재 HPE가 이러한 GPU를 탑재한 서버를 제공하는지는 확인할 수 없었지만, 인텔 데이터 센터 GPU Max 1100 GPU를 탑재한 HPE ProLiant DL380 Gen11 및 DL380a Gen11은 제공한다.²⁷ 즉, Dell PowerEdge XE9640은 현재 2U 서버에 4개의 인텔 Max 1550 OAM GPU를 탑재한 유일한 오퍼링일 수 있다. 데이터 센터 공간과 전력 효율성에 관심이 있는 기업에는 4개의 인텔 Max 1550 GPU를 탑재한 2U 서버가 데이터 센터 공간에는 부담 없이 고성능 컴퓨팅과 에너지 효율성을 결합한 솔루션을 제공한다.

공개된 MLPerf® 3.1 결과에서 4개의 HGX H100 GPU를 탑재한 Dell PowerEdge XE9640은 4개의 PCIe H100 GPU를 탑재한 HPE ProLiant DL380a보다 최대 1.99배 더 우수한 성능을 보였다(그림 3 참조).

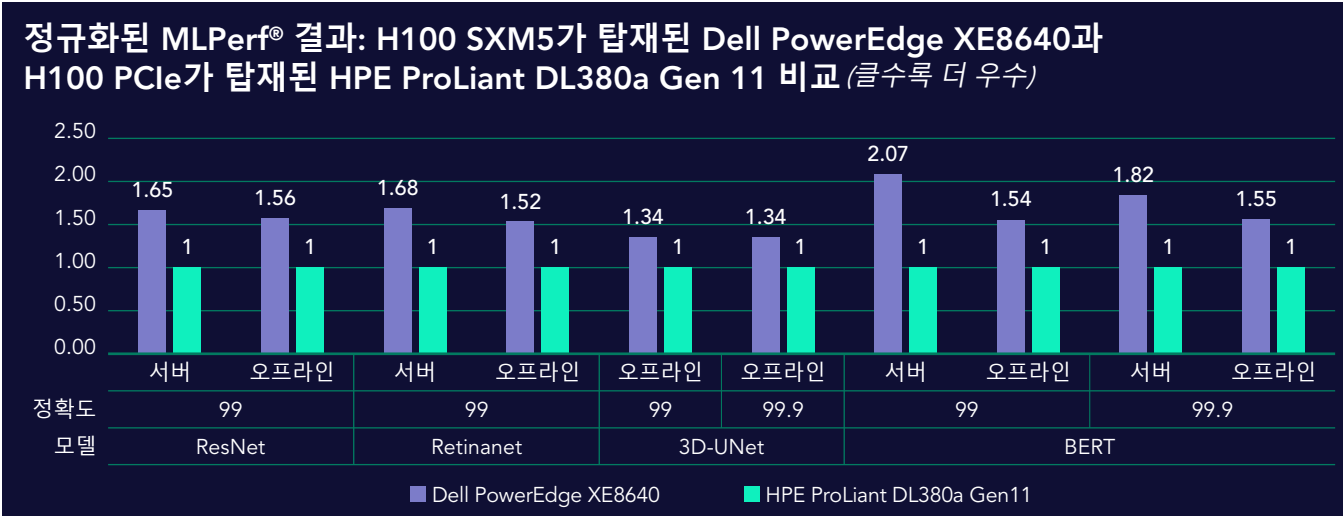


3: 2023년 11월 29일 현재 Dell PowerEdge XE9680 및 HPE ProLiant XL675d Gen10 Plus에 대해 공개된 MLPerf® 결과. Dell 시스템은 NVIDIA H100 GPU를 사용하지만 HPE 시스템의 GPU는 한 세대 이전 버전이다. 출처: MLCommons®의 데이터를 사용하는 Principled Technologies.^{28,29}

PowerEdge XE8640은 프로세서를 위한 공기 냉각과 GPU를 위한 액체 보조 공기 냉각 라디에이터를 갖춘 4웨이 GPU 구성을 제공하므로 랙에 시설 용수를 공급할 필요가 없다. 외부 냉각수를 사용하지 않거나 사용할 수 없는 사용자를 위해³⁰ 4U Dell PowerEdge XE8640은 4개의 NVIDIA H100 SXM5 GPU를 지원하여 직접 액체 냉각 필요 없이 PowerEdge XE9640과 동일한 컴퓨팅 성능을 제공한다.³¹

Dell PowerEdge XE8640은 최신 4세대 인텔 제온 스케일러블 프로세서와 최대 4TB의 메모리를 탑재하여³² AI 및 데이터 분석에서 일반적인 대규모 데이터 세트와 복잡한 계산을 처리한다. 다시 말해, HPE는 HPE Cray 시스템에서 NVIDIA H100 SXM5 GPU를 제공하지만 HPE ProLiant GPU 지원 서버에서는 이를 지원하지 않는다.

2023년 11월에 공개된 MLPerf® 데이터와 비교했을 때 4개의 NVIDIA H100 SXM5 GPU를 탑재한 PowerEdge XE8640 서버는 9가지 범주에서 4개 GPU 제출물 가운데 가장 높은 AI 처리량을 달성했다. 그림 4에서 볼 수 있듯이, HPE ProLiant DL380a 서버에 비해 최대 2.07배 높은 점수를 받았다.

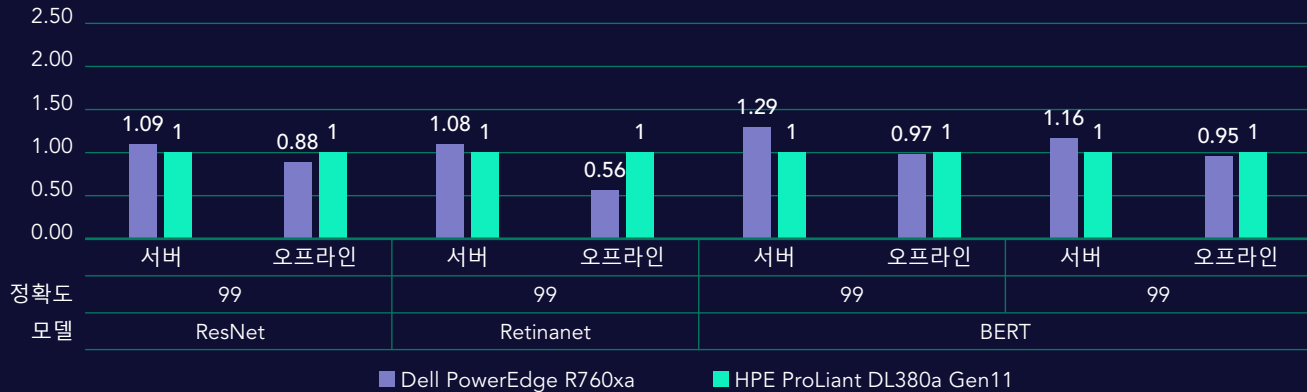


4: 2023년 11월 29일 현재 Dell PowerEdge XE8640 및 HPE ProLiant DL380a Gen11에 대해 공개된 MLPerf® 결과. Dell 시스템은 NVIDIA H100 SXM5 폼 팩터를 사용하고, HPE 시스템은 덜 강력한 PCIe 폼 팩터를 사용한다. 출처: MLCommons®의 데이터를 사용하는 Principled Technologies.^{33,34}

마지막으로, 소규모로 시작하고 필요에 따라 확장하려는 조직의 경우 2U Dell PowerEdge R760xa 서버는 최대 4개의 더블 와이드 PCIe Gen 5 GPU 또는 12개의 싱글 와이드 PCIe GPU를 지원하는 NVIDIA, AMD 및 인텔의 다양한 GPU를 수용한다.³⁵ 32개의 DIMM 슬롯, 2.5" 디스크용 8개 드라이브 베이, 12개의 PCIe 슬롯을 갖추고 있어 증가하는 AI 데이터 요구 사항에 맞춰 확장 가능한 스토리지를 제공하며 NVIDIA H100 또는 L40S와 같은 최대 12개의 싱글 와이드 PCIe GPU 또는 4개의 더블 와이드 PCIe GPU를 지원한다.³⁶ 이러한 확장성은 서버가 머신 러닝 모델 학습에서 고급 데이터 처리에 이르기까지 진화하는 AI 작업에 적응할 수 있음을 의미한다.

PowerEdge R760xa의 공기 냉각 시스템은 고집적 컴퓨팅 환경을 지원하고 최대 350W의 더 높은 TDP(Thermal Design Power) 가속기를 수용할 수 있어³⁷ 집중적인 컴퓨팅 부하에서도 성능을 유지하는 데 도움이 될 수 있다. 2023년 11월 현재 "서버" 모드를 사용하여 공개된 MLPerf® ResNet, RetinaNet 및 BERT 서버 테스트 결과에서 4개의 NVIDIA H100 PCIe GPU를 탑재한 PowerEdge R760xa는 역시 4개의 H100 PCIe GPU를 탑재한 HPE ProLiant DL380a Gen 11보다 우수한 성능을 보였다(그림 5 참조).

정규화된 MLPerf® 결과: H100 PCIe가 탑재된 Dell PowerEdge R760xa와 H100 PCIe가 탑재된 HPE ProLiant DL380a Gen 11 비교 (클수록 더 우수)



5: 2023년 11월 29일 현재 Dell PowerEdge R760xa 및 HPE ProLiant DL380a Gen11에 대해 공개된 MLPerf® 결과. 두 시스템 모두 NVIDIA H100 GPU의 PCIe 폼 팩터를 사용한다. 출처: MLCommons®의 데이터를 사용하는 Principled Technologies.^{38,39}

전반적으로, MLPerf® 결과에 따르면 서버와 구성 요소에 따라 성능이 크게 다르므로 워크로드와 성능 요구 사항을 지원하는 데 적합한 옵션을 선택하는 것이 중요하다. AI 워크로드를 위한 Dell PowerEdge 서버는 강력한 MLPerf® 성능을 제공하면서 기업의 모든 데이터 센터 요구 사항에 맞춰 다양한 냉각 및 집적도 옵션을 제공한다.

Dell AI 포트폴리오에 대한 자세한 내용

컴퓨팅 성능은 중요하지만, AI 워크로드를 계획할 때 고려해야 할 사항 중 하나에 불과하다. AI 구현에 착수할 때 공급업체가 제공하는 나머지 AI 포트폴리오도 고려해야 한다. 아래에서는 클라이언트 워크스테이션, 클라우드 네이티브 제품, 스토리지 등 이러한 AI 포트폴리오에 중요한 추가 범주에 대해 살펴보겠다. 또한 Dell의 오퍼링이 HPE에 비해 어떤 면에서 더 유리한지도 명확히 제시하겠다.

워크스테이션

AI 개발자와 데이터 과학자를 위해 Dell Precision 데이터 사이언스 워크스테이션은 NVIDIA RTX™ GPU 및 인텔 제온® CPU와 데이터 사이언스 툴 제품군을 제공한다.⁴⁰ 이러한 시스템은 100개 이상의 전문 애플리케이션에 대해 인증된 NVIDIA GPU⁴¹와 인텔 제온 스케일러블 프로세서 가속기(예: 인텔 DL 부스트)를 갖춘 전문가급 컴퓨팅 옵션을 활용한다.⁴² Precision 워크스테이션은 모바일, 타워 및 랙 형식으로 제공되어 대규모 고정형 데이터 분석부터 이동식 과학 현장 모델링까지 다양한 요구 사항을 충족한다.

HPE 워크스테이션 오퍼링은 범위가 좁으며, 주로 NVIDIA L4가 장착된 단일 워크스테이션 타워로 구성된다. HPE는 모바일 워크스테이션 옵션을 제공하지 않는다.⁴³ HPE 워크스테이션 타워 오퍼링은 많은 작업에 적합하지만, Dell이 제공하는 것만큼 더 광범위한 유연성과 워크로드 범위를 제공하지는 못한다. Dell Precision 워크스테이션은 크기와 휴대성 옵션이 다양하여 실험실, 사무실 및 현장 작업 등 다양한 환경에서의 요구에 맞는 맞춤형 솔루션을 제공할 수 있다.

스토리지

AI 워크로드를 실행하려면 컴퓨팅만큼이나 스토리지도 중요할 수 있다. 데이터가 많을수록 AI 모델 정확도가 향상되지만, 방대한 데이터 세트를 저장하고 관리하는 것은 많은 데이터 센터의 역량에 부담이 될 수 있다. 또한 모델은 일반적으로 비정형 데이터를 사용하여 학습되기 때문에 AI 지원 스토리지 시스템은 다양한 데이터 유형을 쉽게 처리해야 한다.⁴⁴ AI, ML 및 DL 데이터 세트의 용량과 확장성을 제공하기 위해 Dell은 파일 스토리지용 PowerScale™ Series와 오브젝트 스토리지용 ECS(Elastic Cloud Storage) 또는 소프트웨어 정의 ObjectScale을 제공한다.

PowerScale 올플래시 NAS 포트폴리오는 노드당 3.84TB부터 최대 720TB의 물리적 용량에 이르는 용량 옵션을 제공하며, 클러스터링된 올플래시 용량은 186PB의 물리적 용량에 달한다. PowerScale의 유연성과 확장성은 다양한 고객과 AI 활용 사례를 지원할 수 있다.⁴⁵ PowerScale F900을 클러스터링하면 최대 186PB의 총 물리적 스토리지를 확보할 수 있다.⁴⁶ 세 가지 올플래시 PowerScale 모델(F200, F600, F900) 모두 스토리지 효율성을 개선하기 위해 인라인 데이터 압축 및 중복 제거가 포함되어 있다.⁴⁷ 각 PowerScale 스토리지 모델은 Dell OneFS™ 파일 시스템을 사용하는데, 이 파일 시스템은 정책을 사용하여 스토리지를 계층화하고 워크로드 최적화를 위해 가장 중요한 데이터를 가장 빠른 계층으로 우선순위를 지정한다.⁴⁸ Dell은 또한 AWS 마켓플레이스에서 Dell APEX File Storage for AWS를 통해 OneFS 소프트웨어를 제공한다. 고객은 온프레미스 OneFS 어레이에서 사용 가능한 동일한 기능을 통해 일관된 사용자 경험을 위해 AWS 컴퓨팅 인스턴스와 함께 OneFS를 활용할 수 있다.⁴⁹ HPE는 하이브리드 스토리지 솔루션을 위한 퍼블릭 클라우드 통합을 제공하지만, Dell APEX File Storage for AWS와 같은 클라우드 네이티브 옵션은 발견되지 않았다.

Dell의 오브젝트 스토리지 옵션에는 "퍼블릭 클라우드 규모로 비정형 데이터를 저장하도록 특별히 설계된" Dell ECS(Enterprise Object Storage)가 포함된다.⁵⁰ ECS 스토리지 노드는 하이브리드 클라우드 기능을 위한 Amazon S3 오브젝트 스토리지와의 기본 호환성과 함께 랙당 최대 14PB의 용량을 제공한다.⁵¹ HPE 역시 파일 및 오브젝트 스토리지 옵션을 갖춘 비정형 스토리지를 제공하지만, 오브젝트 스토리지 오퍼링은 Scality와의 파트너십을 통해 제공된다. 고객은 HPE에서 Scality용 HPE 솔루션을 구매할 수 있다.⁵²

전문 서비스

Dell Technologies는 컨설팅, 데이터 준비, 배포, 지원 및 AI 배포를 지원하기 위한 매니지드 서비스를 포함한 광범위한 전문 서비스를 제공한다. 검증된 아키텍처와 솔루션을 찾는 조직을 위해 Dell Technologies는 특정 활용 사례를 대상으로 하여 AI 리소스를 설계하고 배포할 때 추측을 배제하는 Validated Designs for AI를 제공한다. Dell Technologies에서 검증한 이 AI 솔루션에는 하드웨어 및 소프트웨어 번들, 대화형 AI 모델, 머신 러닝 운영 등이 포함된다. Dell Technologies는 사전 구성되고 특별히 설계된 솔루션을 AI 관련 서비스와 결합하여 다양한 AI 요구 사항에 맞는 포괄적인 AI 솔루션을 제공한다. 이러한 오퍼링은 임시 솔루션을 구축하는 것보다 더 빠르고 쉬운 AI 성공 경로를 제시할 수 있다.

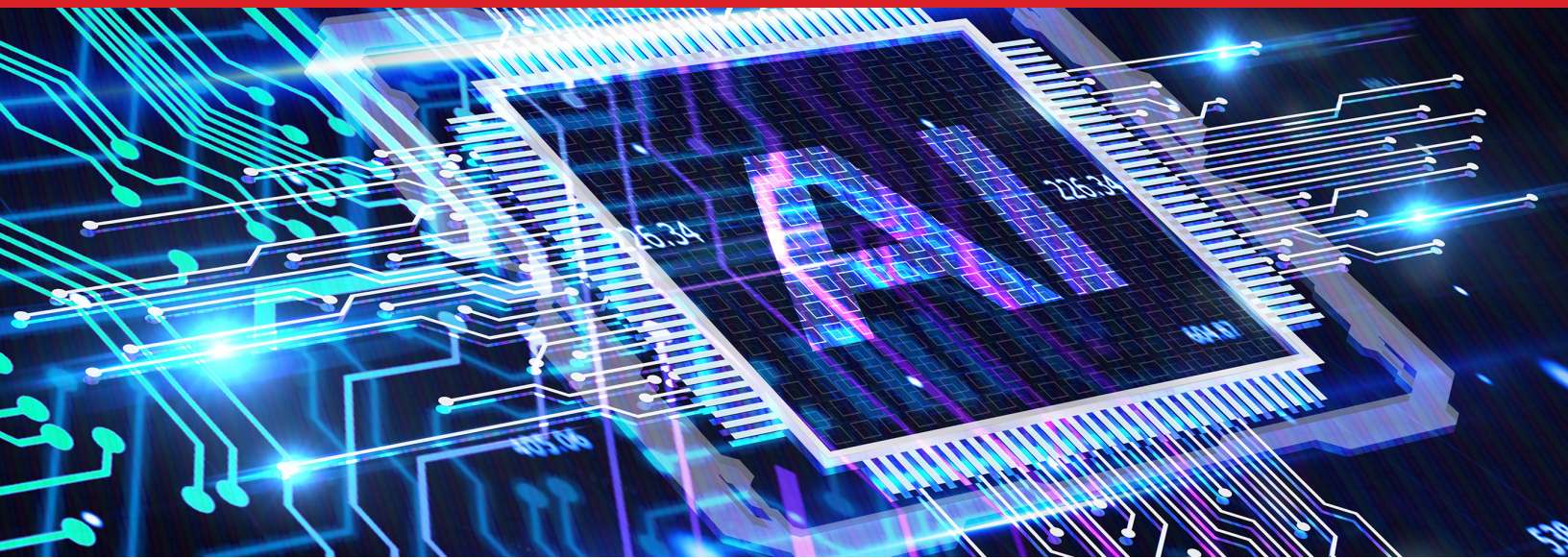
또한 Dell 서비스는 조언에서 구현에 이르는 AI 여정을 안내할 수 있다. Dell ProConsult Advisory Services는 고객이 GenAI 프로세스를 도입하여 어떤 부분에서 이점을 얻을 수 있는지 파악하고 필요한 솔루션과 IT 기술로 구성된 로드맵을 만드는 데 도움이 된다. Dell 서비스는 대규모 언어 모델 통합을 위한 데이터를 준비하고 IT 팀에 AI 지식을 교육할 수 있다. GenAI를 완전히 도입하기 위해 Dell 팀은 특정 활용 사례를 검토하고 요구 사항에 가장 적합한 AI 모델을 결정, 배포 및 구성한다. HPE 역시 기업의 AI 분야 활동을 지원하기 위한 전문 서비스를 제공한다.^{53,54}

관리 고려 사항

서버에는 관리 시간이 소요되는 지속적인 관리가 필요하다. 펌웨어, 소프트웨어 및 드라이버는 주기적인 업데이트가 필요하며, IT 직원이 성능과 온도 등을 최적화하고 유지 관리해야 한다. 이전의 PT(Principled Technologies) 테스트에서는 iDRAC9(Integrated Dell Remote Access Controller 9)을 사용하는 Dell 서버의 관리 기능을 평가했다.⁵⁵ 관리자는 구성 가능한 스케줄링을 갖춘 iDRAC9 OME(OpenManage™ Enterprise)의 자동화된 온라인 업데이트를 통해 서버를 최신 상태로 유지하고, 워크로드가 증가함에 따라 프로파일을 사용하여 빠르고 쉽게 새로운 서버를 탑재할 수 있다. iDRAC 및 OME를 통해 Dell 고객은 HPE OneView 및 HPE iLO를 사용하는 것보다 더 쉽게 더 많은 원격 관리 기능에 액세스하고, 서버를 배포하고, 펌웨어를 업데이트할 수 있다. Dell PowerEdge 서버에는 Dell 관리 및 서비스가 제공되어 조직이 "시스템 상태 모니터링 또는 펌웨어 업데이트와 같은 작업에 소요되는 시간과 노력을 줄여" 혁신과 기타 작업에 필요한 IT 주기를 확보하는 데 도움이 될 수 있다.⁵⁶

3: 2022년 11월 PT 보고서에서 Dell과 HPE 관리 툴 간 비교 요약.⁵⁷
출처: Principled Technologies.

	Dell 관리 툴의 차이점	우수성
더 많은 원격 관리 기능 iDRAC vs. iLO	iDRAC의 더 많은 원격 기능을 위한 더 많은 HTML5 콘솔 및 BIOS 구성 기능	HTML5 콘솔 기능 2.5배, BIOS 기능 13배
더 간편한 서버 구축 OME vs. OneView	OME의 일대다 프로파일 배포	OneView를 사용할 때보다 서버 구축 시간 52% 단축
더 간편한 펌웨어 업데이트 OME vs. OneView	OME의 자동화된 온라인 업데이트	Dell.com에 연결하여 여러 서버를 업데이트하면 OneView를 통해 번들을 수동으로 업로드하여 서버를 업데이트하는 것보다 시간이 절약됨
더 간편한 알림 OME vs. OneView	OME에서 알림 정책을 설정하고 알림을 기반으로 자동화된 작업 실행	이 프로세스를 자동화하면 OneView에서 알림을 수신할 때마다 작업을 수동으로 실행하는 것보다 시간이 절약되고 오류 가능성이 감소함
더 간편한 보안 기능 사용(시스템 잠금 및 동적 USB) iDRAC vs. iLO	iDRAC를 사용하면 재부팅 없이 단계 및 시간 단축	시스템 잠금에 소요되는 단계는 ¼, 시간은 91%로 단축됨
더욱 강력한 분석 PowerEdge용 CloudIQ vs. InfoSight	PowerEdge용 CloudIQ의 맞춤형 보고서, 더 효과적인 관리자 제어를 위한 더 많은 상태 메트릭	InfoSight에 비해 15배 이상 더 많은 메트릭 선택 가능



결론

AI의 능력을 활용해 비즈니스 운영을 간소화하고 개선하는 것은 비즈니스에 상당한 영향을 미치는 까다로운 작업일 수 있다. 기술이 그 어느 때보다 빠르게 발전함에 따라 AI에 적합한 공급업체와 협력하는 것이 중요하다. 포괄적인 AI 포트폴리오를 제공할 뿐만 아니라 계획, 준비, 구현 및 관리 서비스까지 제공할 수 있는 Dell Technologies와 같은 회사를 선택하면 고객은 이러한 당면 과제를 정면으로 해결할 수 있다. MLPerf® 벤치마크 테스트 결과에 따르면 Dell AI 포트폴리오의 오퍼링은 AI 워크로드에 일관된 강력한 성능을 제공한다. Dell은 고성능의 유연한 서버 옵션과 다양한 스토리지 선택, 검증된 솔루션, AI에 특별히 맞춤화된 전문 서비스를 통해 기업이 AI와 그 이점을 수용하도록 지원할 수 있다.

1. Dell, "Increasing Your Data Value with Dell Generative AI Solutions"(2023년 12월 19일 액세스), <https://www.dell.com/en-us/blog/increasing-your-data-value-with-dell-generative-ai-solutions/>.
2. Dell, "Dell AI 솔루션"(2023년 12월 12일 액세스), <https://www.dell.com/en-us/dt/solutions/artificial-intelligence/index.htm#accordion0&tab0=0>.
3. Dell, "Dell Technologies and Hugging Face to Simplify Generative AI with On-Premises IT"(2023년 12월 12일 액세스), <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#/filter-on/Country:en-us>.
4. Dell, "Dell and Meta Collaborate to Drive Generative AI Innovation"(2023년 12월 12일 액세스), <https://www.dell.com/en-us/blog/dell-and-meta-collaborate-to-drive-generative-ai-innovation/>.
5. Dell, "Dell AI 지원 데이터 플랫폼"(2023년 12월 12일 액세스), <https://www.dell.com/en-us/dt/solutions/artificial-intelligence/storage-for-ai.htm?hve=explore+unstructured+storage#tab0=0>.
6. Dell, "Snowflake and Dell Partnership Gains Momentum"(2023년 12월 19일 액세스), <https://www.dell.com/en-us/blog/snowflake-and-dell-partnership-gains-momentum/>.
7. Robert McNeal, "Dell, VMware and NVIDIA Bring AI to Your Data"(2024년 1월 17일 액세스), <https://www.dell.com/en-us/blog/dell-vmware-and-nvidia-bring-ai-to-your-data/>. 위 링크에 따름: "Dell 분석 기준, 2023년 8월, Dell Technologies는 워크스테이션 PC(모바일 및 고정형)부터 HPC(High Performance Computing) 서버, 데이터 스토리지, 클라우드 네이티브 소프트웨어 정의 인프라스트럭처, 네트워킹 스위치, 데이터 보호, HCI 및 서비스에 이르기까지 AI 워크로드를 지원하도록 엔지니어링된 솔루션을 제공한다."
8. 인텔, "인텔 Advanced Matrix Extensions로 인공 지능 워크로드 가속화(인텔 AMX)"(2023년 12월 12일 액세스), <https://www.intel.com/content/www/us/en/content-details/785250/accelerate-artificial-intelligence-ai-workloads-with-intel-advanced-matrix-extensions-intel-amx.html>.
9. Vipera, "NVIDIA's H100 and A100 GPU Cards: Exploring the Intricacies of SXM and PCI-E Connections,"(2023년 12월 12일 액세스), <https://www.viperatech.com/unraveling-the-mysteries-sxm-vs-pci-e-connections-in-nvidias-high-end-h100-and-a100-gpus/>.

10. GitHub, "MLPerf® Results Messaging Guidelines"(2024년 1월 16일 액세스), https://github.com/mlcommons/policies/blob/master/MLPerf_Results_Messaging_Guidelines.adoc.
11. MLCommons®, "MLPerf® Inference: Datacenter Benchmark Suite Results"(2023년 12월 12일 액세스), <https://mlcommons.org/en/inference-datacenter-31/>.
12. Dell, "PowerEdge XE 서버"(2023년 12월 12일 액세스), <https://www.dell.com/en-us/dt/servers/specialty-servers/poweredge-xe-servers.htm?hve=explore+poweredge+xe#tab0=0>.
13. HPE, "HPE ProLiant XL675d Gen10 Plus Configure-to-order Server"(2023년 12월 12일 액세스), <https://www.hpe.com/us/en/product-catalog/compute/proliant-servers/pip.1013142988.html>.
14. HPE, "HPE ProLiant DL380a Gen11"(2023년 12월 12일 액세스), <https://www.hpe.com/us/en/product-catalog/compute/proliant-servers/pip.proliant-dl380-server.1014696168.html>.
15. MLCommons®, "MLPerf® Inference: Datacenter Benchmark Suite Results v 3.1"(2023년 12월 12일 액세스), <https://mlcommons.org/benchmarks/inference-datacenter/>.
16. HPE, "NVIDIA Accelerators for HPE ProLiant Servers"(2023년 12월 12일 액세스), https://www.hpe.com/psnow/doc/c04123180.html?jumpid=in_pdp-psnow-qs.
17. Dell, "PowerEdge XE9680 스펙 시트"(2024년 1월 19일 액세스), <https://www.delltechnologies.com/asset/en-us/products/servers/technical-support/poweredge-xe9680-spec-sheet.pdf>.
18. HPE, "HPE & NVIDIA financial services solution sets new records in performance"(2023년 12월 12일 액세스), <https://community.hpe.com/t5/alliances/hpe-amp-nvidia-financial-services-solution-sets-new-records-in/ba-p/7197388>.
19. HPE, "QuickSpecs: HPE Cray Supercomputing XD670"(2023년 12월 12일 액세스), <https://www.hpe.com/psnow/doc/a50004292enw>.
20. 검증된 MLPerf® v3.1 Inference Closed 점수. <https://mlcommons.org/benchmarks/inference-datacenter/> 2023년 12월 5일, 3.1-0069 항목에서 발췌. MLPerf® 이름 및 로고는 미국 및 기타 국가에서 MLCommons® Association의 등록 및 미등록 상표입니다. All rights reserved. 무단 사용을 엄격히 금지합니다. 자세한 내용은 www.mlcommons.org 를 참조하십시오.
21. 검증된 MLPerf® v3.1 Inference Closed 점수. <https://mlcommons.org/benchmarks/inference-datacenter/> 2023년 12월 5일, 3.1-0085 항목에서 발췌. MLPerf® 이름 및 로고는 미국 및 기타 국가에서 MLCommons® Association의 등록 및 미등록 상표입니다. All rights reserved. 무단 사용을 엄격히 금지합니다. 자세한 내용은 www.mlcommons.org 를 참조하십시오.
22. Dell, "PowerEdge XE9640 랙 서버"(2023년 12월 12일 액세스), <https://www.dell.com/en-us/shop/ipovw/poweredge-xe9640>.
23. Accelsius, "Enabling the AI Revolution with Liquid Cooling"(2023년 12월 12일 액세스), <https://www.accelsius.com/blog/enabling-the-ai-revolution-with-liquid-cooling>.
24. Dell, "Dell PowerEdge XE9640 기술 가이드"(2023년 12월 12일 액세스), <https://www.delltechnologies.com/asset/en-us/products/servers/technical-support/poweredge-xe9640-technical-guide.pdf>.
25. HPE, "HPE ProLiant DL380a Gen11"(2023년 12월 12일 액세스), https://www.hpe.com/psnow/doc/PSN1014696168WWEN.pdf?jumpid=in_pdp-psnow-dds.
26. 인텔, "인텔® Data Center GPU Max Series 기술 개요"(2023년 12월 12일 액세스), <https://www.intel.com/content/www/us/en/developer/articles/technical/intel-data-center-gpu-max-series-overview.html#gs.08874l>.
27. HPE, "Intel Data Center GPU Max 1100 48GB Accelerator for HPE Data sheet"(2023년 12월 12일 액세스), <https://www.hpe.com/psnow/doc/PSN1014779728WWEN>.
28. 검증된 MLPerf® v3.1 Inference Closed 점수. <https://mlcommons.org/benchmarks/inference-datacenter/> 2023년 12월 5일, 3.1-0066 항목에서 발췌. MLPerf® 이름 및 로고는 미국 및 기타 국가에서 MLCommons® Association의 등록 및 미등록 상표입니다. All rights reserved. 무단 사용을 엄격히 금지합니다. 자세한 내용은 www.mlcommons.org 를 참조하십시오.
29. 검증된 MLPerf® v3.1 Inference Closed 점수. <https://mlcommons.org/benchmarks/inference-datacenter/> 2023년 12월 5일, 3.1-0084 항목에서 발췌. MLPerf® 이름 및 로고는 미국 및 기타 국가에서 MLCommons® Association의 등록 및 미등록 상표입니다. All rights reserved. 무단 사용을 엄격히 금지합니다. 자세한 내용은 www.mlcommons.org 를 참조하십시오.
30. Dell, "AI and HPC —With Air or Liquid Cooling"(2023년 12월 12일 액세스), <https://www.delltechnologies.com/asset/en-us/products/servers/briefs-summaries/poweredge-xe9640-and-xe8640-infographic.pdf>.
31. Dell, "PowerEdge XE8640: AI, HPC 모델링 및 시뮬레이션 워크로드를 사용하여 탁월한 성능 발휘"(2023년 12월 12일 액세스), <https://www.delltechnologies.com/asset/en-us/products/servers/technical-support/poweredge-xe8640-spec-sheet.pdf>.
32. Dell, "PowerEdge XE8640 랙 서버"(2023년 12월 12일 액세스), <https://www.dell.com/en-us/shop/ipovw/poweredge-xe8640>.

-
33. 검증된 MLPerf® v3.1 Inference Closed 점수. <https://mlcommons.org/benchmarks/inference-datacenter/> 2023년 12월 5일, 3.1-0067 항목에서 발췌. MLPerf® 이름 및 로고는 미국 및 기타 국가에서 MLCommons® Association의 등록 및 미등록 상표입니다. All rights reserved. 무단 사용을 엄격히 금지합니다. 자세한 내용은 www.mlcommons.org 를 참조하십시오.
 34. 검증된 MLPerf® v3.1 Inference Closed 점수. <https://mlcommons.org/benchmarks/inference-datacenter/> 2023년 12월 5일, 3.1-0064 항목에서 발췌. MLPerf® 이름 및 로고는 미국 및 기타 국가에서 MLCommons® Association의 등록 및 미등록 상표입니다. All rights reserved. 무단 사용을 엄격히 금지합니다. 자세한 내용은 www.mlcommons.org 를 참조하십시오.
 35. Dell, "PowerEdge R760xa 랙 서버"(2023년 12월 12일 액세스), https://www.dell.com/en-us/shop/dell-poweredge-servers/poweredge-r760xa-rack-server/spd/poweredge-r760xa/pe_r760xa_16902_vi_vp#features_section.
 36. SANStorageWorks, "Dell EMC PowerEdge R760xa: Powerful and scalable for GPU workloads"(2023년 12월 12일 액세스), <https://www.sanstorageworks.com/PowerEdge-R760xa.asp>.
 37. Dell, "Dell PowerEdge Servers and NVIDIA GPUs"(2023년 12월 12일 액세스), <https://infohub.delltechnologies.com//design-guide-generative-ai-in-the-enterprise-inferencing/dell-poweredge-servers-and-nvidia-gpus-1/>.
 38. 검증된 MLPerf® v3.1 Inference Closed 점수. <https://mlcommons.org/benchmarks/inference-datacenter/> 2023년 12월 5일, 3.1-0064 항목에서 발췌. MLPerf® 이름 및 로고는 미국 및 기타 국가에서 MLCommons® Association의 등록 및 미등록 상표입니다. All rights reserved. 무단 사용을 엄격히 금지합니다. 자세한 내용은 www.mlcommons.org 를 참조하십시오.
 39. 검증된 MLPerf® v3.1 Inference Closed 점수. <https://mlcommons.org/benchmarks/inference-datacenter/> 2023년 12월 5일, 3.1-0064 항목에서 발췌. MLPerf® 이름 및 로고는 미국 및 기타 국가에서 MLCommons® Association의 등록 및 미등록 상표입니다. All rights reserved. 무단 사용을 엄격히 금지합니다. 자세한 내용은 www.mlcommons.org 를 참조하십시오.
 40. Dell, "AI를 위한 워크스테이션"(2023년 12월 12일 액세스), <https://www.dell.com/en-us/dt/ai-technologies/index.htm?hve=explore+dell+precision+for+ai#pdf-overlay=//www.delltechnologies.com/asset/en-us/products/workstations/briefs-summaries/ai-industry-brochure.pdf>.
 41. NVIDIA, "NVIDIA RTX in Professional Workstations"(2023년 12월 12일 액세스), <https://www.nvidia.com/en-us/design-visualization/desktop-graphics/>.
 42. 인텔, "Intel® Deep Learning Boost (Intel® DL Boost)"(2023년 12월 12일 액세스), <https://www.intel.com/content/www/us/en/artificial-intelligence/deep-learning-boost.html>.
 43. HPE, "HPE ProLiant ML350 Gen11"(2023년 12월 12일 액세스), <https://buy.hpe.com/us/en/compute/tower-servers/proliant-ml300-servers/proliant-ml350-server/hpe-proliant-ml350-gen11/p/1014696172>.
 44. ComputerWeekly.com, "Storage requirements for AI, ML and analytics in 2022"(2023년 12월 12일 액세스), <https://www.computerweekly.com/feature/Storage-requirements-for-AI-ML-and-analytics-in-2022>.
 45. Dell, "PowerScale AI 지원 데이터 플랫폼"(2023년 12월 12일 액세스), <https://www.dell.com/en-us/shop/powerscale-family/sf/powerscale>.
 46. Dell, "PowerScale 비교"(2023년 12월 12일 액세스), <https://www.dell.com/en-us/shop/powerscale-family/sf/powerscale#compare-module>.
 47. Dell, "Dell PowerScale All-Flash"(2023년 12월 12일 액세스), <https://www.delltechnologies.com/asset/en-us/products/storage/technical-support/h15963-ss-powerscale-all-flash-nodes.pdf>.
 48. Dell, "Dell PowerScale OneFS 소프트웨어 기능"(2023년 12월 12일 액세스), <https://www.delltechnologies.com/asset/en-us/products/storage/technical-support/h18275-onefs-software-features-data-sheet.pdf>.
 49. Dell, "Dell APEX File Storage for AWS"(2023년 12월 12일 액세스), <https://www.delltechnologies.com/asset/en-us/products/storage/briefs-summaries/h19575-so-apex-file-storage-for-aws.pdf>.
 50. Dell, "Dell ECS 엔터프라이즈 오브젝트 스토리지"(2023년 12월 12일 액세스), <https://www.dell.com/en-us/dt/storage/ecs/index.htm?hve=explore+ecs#tab0=0&tab1=0>.
 51. Dell, "Dell ECS 엔터프라이즈 오브젝트 스토리지"(2023년 12월 12일 액세스), <https://www.dell.com/en-us/dt/storage/ecs/index.htm#tab0=0&tab1=0&accordion0>.
 52. HPE, "Storage Solutions for Scalify"(2023년 12월 12일 액세스), <https://www.hpe.com/us/en/storage/file-object/scalify.html>.
 53. HPE, "Make AI work for you"(2024년 1월 16일 액세스), <https://www.hpe.com/us/en/solutions/ai-artificial-intelligence.html>.
 54. HPE, "HPE AI Services – Generative AI Implementation"(2024년 1월 16일 액세스), <https://www.hpe.com/us/en/services/generative-ai-implementation-service.html>.

55. Principled Technologies, "Dell 관리 포트폴리오의 툴과 HPE의 유사 툴을 비교 사용하여 관리자 작업을 단순화하고 보안 및 상태 모니터링을 개선"(2023년 12월 12일 액세스),
<https://www.principledtechnologies.com/Dell/Management-tools-vs-HPE-1122.pdf>.
56. Principled Technologies, "Dell 관리 포트폴리오의 툴과 HPE의 유사 툴을 비교 사용하여 관리자 작업을 단순화하고 보안 및 상태 모니터링을 개선"(2023년 12월 12일 액세스),
<https://www.principledtechnologies.com/Dell/Management-tools-vs-HPE-1122.pdf>.
57. Principled Technologies, "Dell 관리 포트폴리오의 툴과 HPE의 유사 툴을 비교 사용하여 관리자 작업을 단순화하고 보안 및 상태 모니터링을 개선"

MLPerf 이름 및 로고는 미국 및 기타 국가에서 MLCommons Association의 등록 및 미등록 상표입니다. All rights reserved. 무단 사용을 엄격히 금지합니다. 자세한 내용은 www.mlcommons.org를 참조하십시오.

▶ 이 보고서의 원본 영어 버전 확인:
<https://facts.pt/zPmSx4c>

이 프로젝트는 Dell Technologies의 의뢰로 진행되었습니다.



Facts matter.®

Principled Technologies는 Principled Technologies, Inc.의 등록 상표입니다.
다른 모든 제품 이름은 해당 소유주의 상표입니다.

보증 및 책임의 면책:

Principled Technologies, Inc.는 테스트의 정확성과 유효성을 보장하기 위한 합리적인 노력을 기울였습니다. 하지만 Principled Technologies, Inc.는 특정 목적에의 적합성에 대한 암시적인 보증을 비롯하여 테스트 결과 및 분석, 정확성, 완전성 또는 품질에 대한 어떠한 명시적 또는 암시적 보증도 명확히 거부합니다. 테스트 결과 이용에 따른 모든 위험은 해당 개인 또는 단체가 스스로 감수해야 하며 Principled Technologies, Inc., 그 직원 및 하청업체가 테스트 절차 또는 결과의 오류 또는 결함으로 인해 발생하는 손해 또는 배상 소송에 대해 어떠한 책임도 지지 않는다는 것에 동의합니다.

어떠한 경우에도 Principled Technologies, Inc.는 테스트와 관련하여 발생하는 간접적, 특별적, 우발적 또는 파생적 손해에 대해 책임지지 않습니다. 이는 그러한 손해의 가능성을 사전에 인지한 경우에도 마찬가지입니다. 어떠한 경우에도 직접적 손해를 포함한 Principled Technologies, Inc.의 책임은 Principled Technologies, Inc.에게 지불된 테스트 비용을 초과하지 않습니다. 여기에 명시된 내용이 고객의 유일한 구매계약입니다.