



Dell AI 포트폴리오를 통해 AI 성공으로 향하는 경로 찾기

Dell AI 포트폴리오와 Supermicro의 유사 오퍼링 비교

AI(Artificial Intelligence)는 산업 전반에 걸쳐 비즈니스 운영을 재편할 준비가 된 새로운 기술 분야이다. 다양한 분야의 조직들이 AI를 활용하여 비즈니스 운영을 개선하는 방법을 모색하고 있지만, AI를 구현하고 그 이점을 누리는 것이 하루아침에 이루어지는 일은 아니라는 점을 기억해야 한다. 모든 비즈니스는 고유하기 때문에 각각의 조직은 데이터와 비즈니스 목표를 평가하여 데이터와 AI를 사용하는 것이 원하는 구체적인 결과를 얻은 데 어떻게 도움이 될 것인지를 확인해야 한다. 계획, 데이터 준비, 적절한 하드웨어 선택, AI 모델 설계, POC(Proof of Concept) 테스트, 레퍼런스 아키텍처, 포괄적인 지원을 비롯한 종합적인 AI 포트폴리오를 제공하는 Dell과 같은 기업과 협력하면 성공적인 AI 프로젝트로 이어질 수 있다.

시장에는 무수히 많은 옵션이 제공되고 있으므로, 이러한 모든 결정을 지원할 수 있는 파트너를 찾는 것이 성공적인 AI 구현과 값비싼 실수 간의 차이를 만들어 낼 수 있다. 이 백서에서는 Dell과 Supermicro의 AI 포트폴리오를 자세히 살펴보면서 고객의 AI 여정 전반에 걸쳐 Dell이 제공할 수 있는 이점을 독자들에게 설명하고자 한다. 먼저 서버 및 컴퓨팅 옵션에 초점을 맞춰 살펴볼 것이다. 각 회사가 고객에게 다양한 오퍼링을 제공하고 있다. 다음으로, Dell이 교육, 계획 서비스, 파트너 생태계 등을 원하는 기업을 위해 하드웨어 외의 것들을 어떻게 고려하고 있는지 살펴볼 것이다.

AI 워크로드에 대한 서버 및 성능 결과

서버는 AI 워크로드를 지원하는 기본 컴퓨팅 인프라스트럭처로서, 워크로드의 크기나 유형에 따라 CPU, GPU 또는 이 두 가지 모두를 컴퓨팅 리소스로 사용할 수 있다. HPC 또는 AI와 같이 더 크고 까다로운 워크로드의 경우 GPU가 최고 수준의 성능을 제공한다. GPU는 범용 PCIe, OAM(Open Compute Project Accelerator Module), 독점 NVIDIA SXM 아키텍처 등 다양한 폼 팩터로 제공되며 현재 최고 수준의 성능을 제공한다.¹ 큰 메모리 용량과 발열 억제 아키텍처 및 전력 효율성과 같은 서버 설계 기능도 성능에 영향을 미친다. 대부분의 데이터 센터는 여전히 공기 냉각을 사용하므로, AI 워크로드에는 가능한 한 효과적으로 공기를 이용해 냉각하도록 제작된 서버가 필요하다. 아래에서는 공개된 MLCommons® MLPerf® 점수와 더불어 구성 요소, 냉각 옵션 등의 측면에서 PowerEdge 서버 오퍼링을 중점적으로 살펴본다.

검사 결과

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.² 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).

표 1: MLPerf® 테스트에 포함된 AI, ML 및 DL 모델 그리고 각 모델에 대한 테스트 및 일반 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 정보

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 모델의 학습 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 11월에 MLCommons® 웹사이트에 게시된 MLPerf® v3.1 Inference Datacenter 결과에서 발췌한 것이다.³ 이러한 결과에는 기술 제조업체와 클라우드 서비스 공급업체의 제출물이 포함되며 다양한 구성이 포괄되어 있다. Supermicro에서 공개적으로 제출한 자료와 비교했을 때 Dell PowerEdge 서버는 비슷한 결과를 얻었다. 표 2에 서버 세부 정보가 나와 있다.

표 2: 2023년 11월 현재 발표된 MLCommons® MLPerf® 3.1 결과에 포함된 Dell 및 Supermicro 서버. 출처: Principled Technologies.

제출자	서버 모델	GPU 수 및 모델	설명
Dell ⁴	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 앱 구동용
Supermicro ⁵	AS-8125GS-TNHR	8x NVIDIA H100 SXM	AMD 프로세서를 사용하는 대규모 AI 학습 및 HPC 워크로드용
	SYS-821GE-TNHR	8x NVIDIA H100 SXM	인텔 프로세서를 사용하는 대규모 AI 학습 및 HPC 워크로드용
	SYS-421GU-TNHR	4x NVIDIA H100 SXM	HPC 및 AI 워크로드를 지원하는 유연성을 위한 모듈형 설계

Dell과 Supermicro가 동일한 GPU 구성에 대한 결과를 제출했으므로 성능을 쉽게 비교할 수 있다. 그림 1과 2에서는 두 공급업체의 동일한 구성이 대체적으로 서로 비슷한 결과로 이어지는 것을 볼 수 있다. 그림 3 및 4에 나온 것과 같은 서로 다른 구성의 경우, 4-GPU 테스트 시 gptj-99.9 모델에서 Dell이 Supermicro를 능가하였다. Dell은 세 가지 서버 모두에서 사용 가능한 모든 모델에 대한 결과를 제출했지만 Supermicro는 그렇지 않았다. 따라서 두 서버 모두에 대해 결과가 있는 모델만 비교한다. 전체 Dell 결과를 보려면 MLCommons® MLPerf® 결과를 참조하면 된다.

8-GPU 서버 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. Dell PowerEdge XE9680은 AMD Instinct™ MI300X 가속기도 제공한다.⁶ 6U의 랙 공간만 차지하는 PowerEdge XE9680은 컴팩트한 8웨이 NVIDIA H100 SXM5 서버이다.

반면, 8-GPU Supermicro 서버는 8U에서 33% 더 많은 랙 공간이 필요하며 NVIDIA GPU를 위한 SXM 폼 팩터도 제공한다. 이러한 크기 차이는 7대의 Dell PowerEdge XE9680 서버를 랙에 장착할 수 있다는 것을 의미한다. Supermicro 서버는 5대만 장착할 수 있다. 그림 1 및 2에서는 Dell PowerEdge XE9680 결과를 Supermicro 8-GPU 서버의 두 구성, 즉 인텔 프로세서가 탑재된 SYS-821GE-TNHR과 AMD 프로세서가 탑재된 AS-8125GS-TNHR의 결과와 비교하고 있다. Supermicro는 SYS-821GE-TNHR에서 RNN-T에 대한 결과를 제출하지 않았으므로 그림 1의 차트에서 이 모델을 제외되었다.

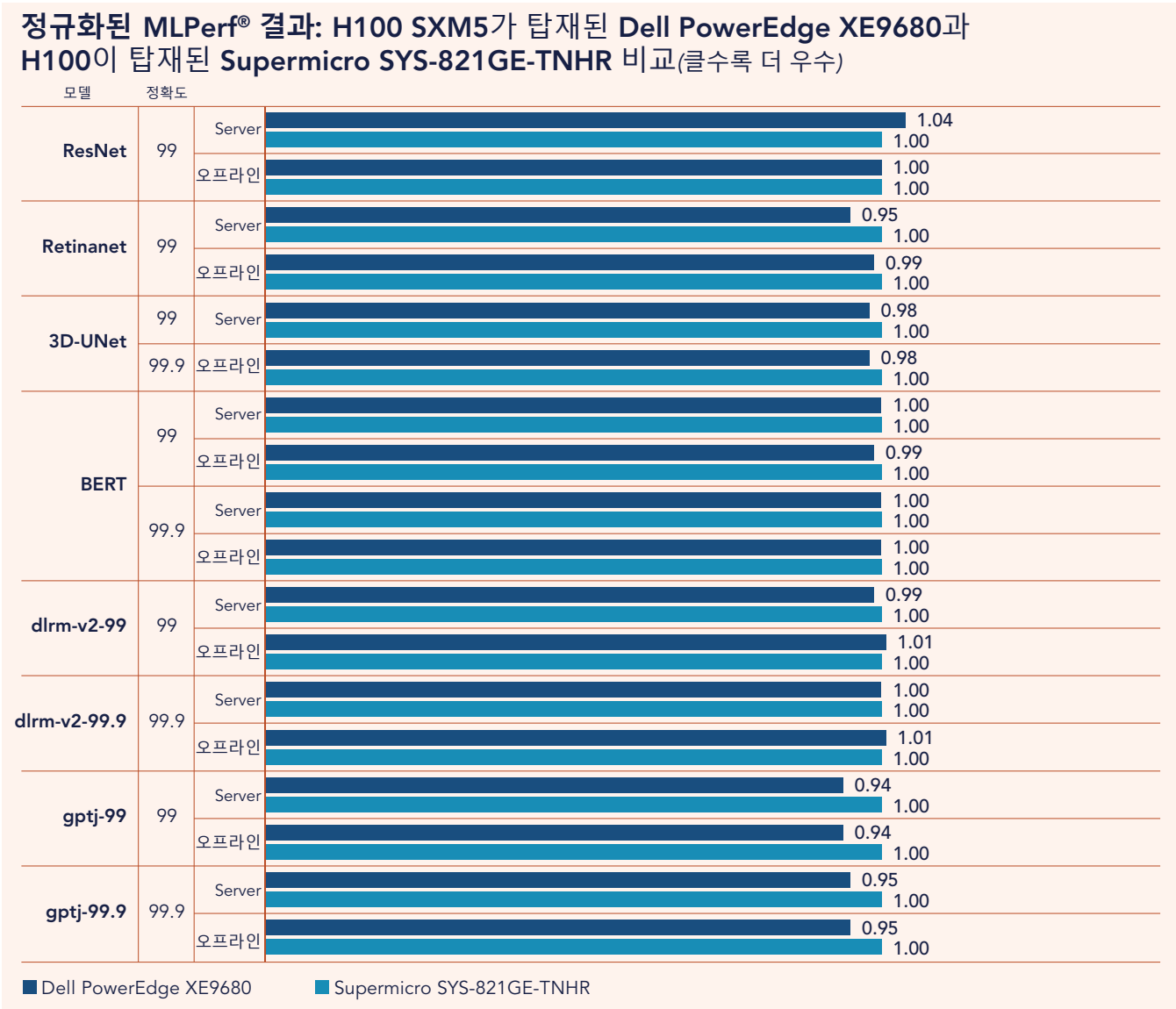


그림 1: 2023년 11월 29일 현재 Dell PowerEdge XE9680 및 Supermicro SYS-821GE-TNHR에 대해 공개된 MLPerf® 결과. 두 시스템 모두 NVIDIA H100 GPU의 SXM 폼 팩터를 탑재하고 있다. 출처: MLCommons®의 데이터를 사용하는 Principled Technologies.^{7,8}



정규화된 MLPerf® 결과: H100 SXM5가 탑재된 Dell PowerEdge XE9680과 H100 SXM5가 탑재된 Supermicro AS-8125GS-TNHR 비교(클수록 더 우수)

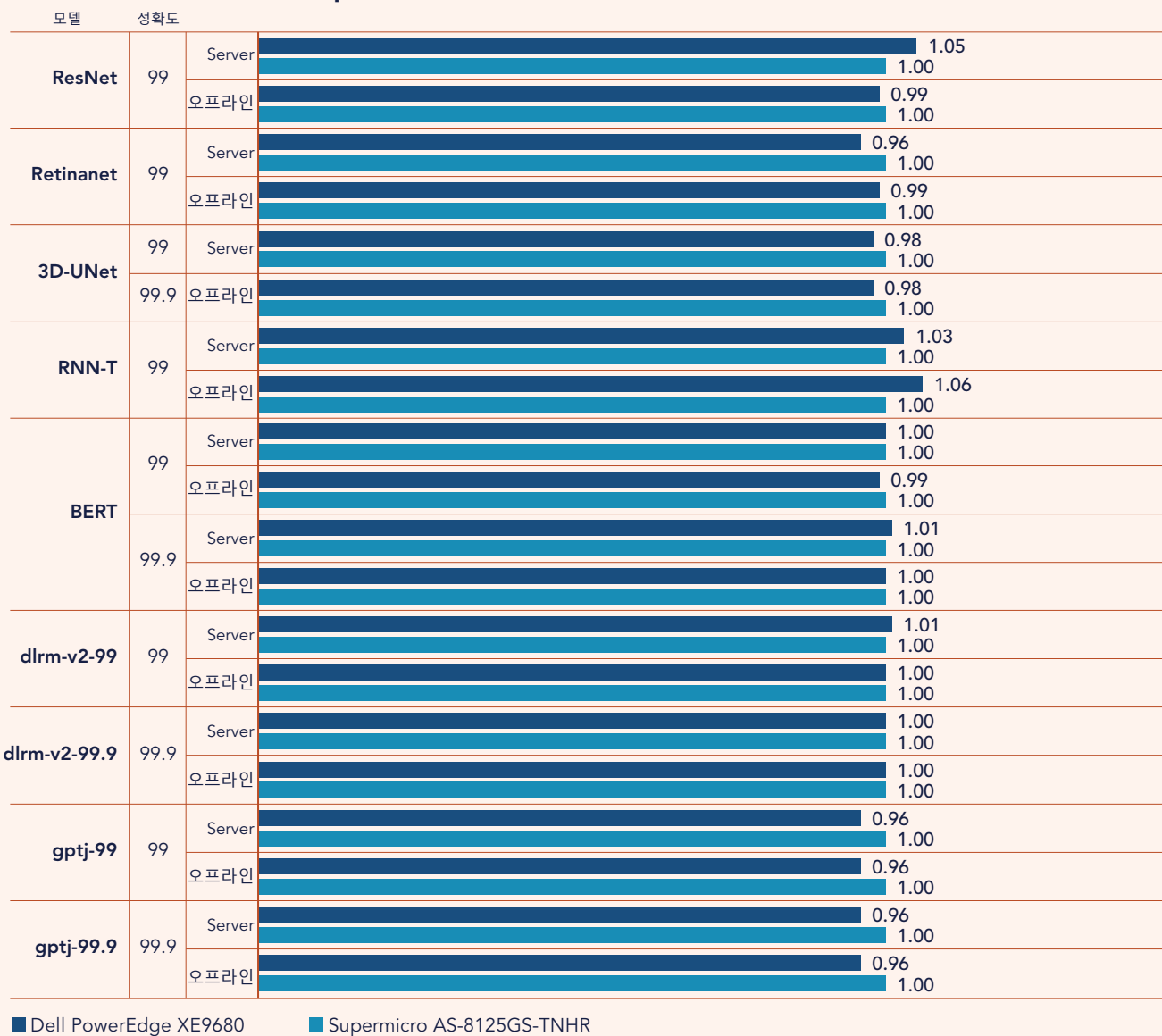


그림 2: 2023년 11월 29일 현재 Dell PowerEdge XE9680 및 Supermicro AS-8125GS-TNHR에 대해 공개된 MLPerf® 결과. 두 시스템 모두 NVIDIA H100 GPU의 SXM 품 팩터를 탑재하고 있다. 출처: MLCommons®의 데이터를 사용하는 Principled Technologies.^{9,10}

이러한 결과에서 알 수 있듯이 Dell PowerEdge XE9680을 선택하면 데이터 센터 공간을 덜 차지하면서 수많은 AI 추론 워크로드에서 유사한 성능을 얻을 수 있다.

4-GPU 서버 결과

데이터 센터 전력 사용량 또는 공간을 최소화하는 것이 주요 관심사라면 2U Dell PowerEdge XE9640이 답이 될 수 있다. 최대 4개의 NVIDIA H100 SXM GPU를 탑재한 PowerEdge XE9640은 PowerEdge XE9680의 절반에 달하는 GPU 컴퓨팅 성능을 1/3의 공간에서 제공한다.¹¹ 밀집된 Dell PowerEdge XE9640 서버는 Dell Smart Cooling 기술을 통합하여 CPU 및 GPU를 위한 직접 액체 냉각 등 다양한 방열 기술을 제공한다.¹² 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.¹³

Supermicro는 2U 폼 팩터에서 4개의 NVIDIA A100 HGX GPU를 제공하는 구형 SYS-220GQ-TNAR+ 서버를 보유하고 있지만, PowerEdge XE9640처럼 최신 H100 HGX GPU 4개가 탑재된 2U Supermicro 서버는 찾을 수 없었다.¹⁴ 제출된 MLPerf® 결과에서 Supermicro의 HGX H100 NVIDIA GPU 탑재 4-GPU 서버는 4U 서버인 SYS-421GU-TNXR이다. 앞서 언급했듯이 Supermicro는 gptj-99.9 AI 모델의 SYS-421GU-TNXR에 대해서만 MLPerf® 3.1 결과를 제출했기 때문에 다른 모델에서는 PowerEdge XE9640과 비교할 수 없다. 그러나 공개된 결과에서는 PowerEdge XE9640이 오프라인 테스트에서 Supermicro 서버를 능가하여 최대 1.37배의 점수를 얻었다(그림 3 참조).

정규화된 MLPerf® 결과: H100 SXM5가 탑재된 Dell PowerEdge XE9640과 H100 SXM5가 탑재된 Supermicro SYS-421GU-TNXR 비교(클수록 더 우수)

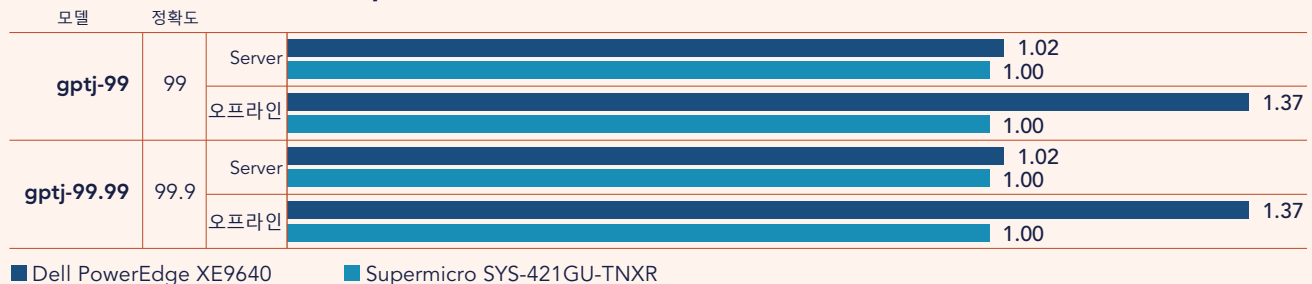


그림 3: 2023년 11월 29일 현재 Dell PowerEdge XE9640 및 Supermicro SYS-421GU-TNXR에 대해 공개된 MLPerf® 결과. 두 시스템 모두 NVIDIA H100 GPU의 SXM 폼 팩터를 탑재하고 있다. 출처: MLCommons®의 데이터를 사용하는 Principled Technologies.^{15,16}

Supermicro SYS-421GU-TNXR 서버 결과를 NVIDIA H100 HGX GPU도 지원하는 4U 4-GPU 서버인 Dell PowerEdge XE8640과 비교할 수도 있다. PowerEdge XE8640은 PowerEdge XE9640보다 크지만 직접 액체 냉각이 필요하지 않으므로, 액체 냉각을 이용할 수 없는 데이터 센터에서 밀도와 냉각 기술 간의 절충이 필요할 때 유용하다. PowerEdge XE8640은 프로세서를 위한 공기 냉각과 GPU를 위한 액체 보조 공기 냉각 라디에이터를 갖추고 있으므로 랙에 시설 용수를 공급해야 할 필요가 없다.¹⁷ Dell PowerEdge XE8640은 최신 4세대 인텔 제온 스케일러블 프로세서와 최대 4TB의 메모리를 탑재하여 AI 및 데이터 분석에서 흔한 대규모 데이터 세트와 복잡한 계산을 처리한다.¹⁸ 4U 폼 팩터의 PowerEdge XE8640은 밀도와 GPU 기능 양쪽 측면 모두에서 Supermicro SYS-421GU-TNXR과 유사하다. 그러나 PowerEdge XE9640에서 보았듯이 Dell PowerEdge XE8640은 오프라인 테스트에서 Supermicro 서버보다 gptj-99 점수가 더 높다(그림 4).

정규화된 MLPerf® 결과: H100 SXM5가 탑재된 Dell PowerEdge XE8640과 H100 SXM5가 탑재된 Supermicro SYS-421GU-TNXR 비교(클수록 더 우수)

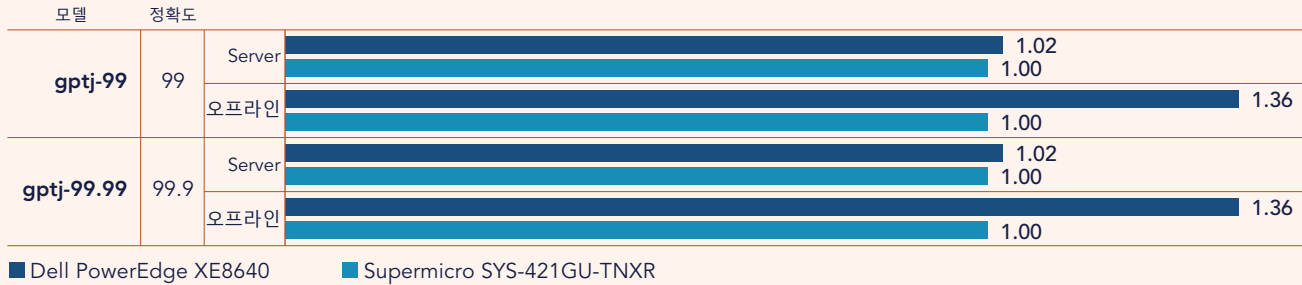


그림 4: 2023년 11월 29일 현재 Dell PowerEdge XE8640 및 Supermicro SYS-421GU-TNXR에 대해 공개된 MLPerf® 결과. 두 시스템 모두 NVIDIA H100 GPU의 SXM 폼 팩터를 탑재하고 있다. 출처: MLCommons®의 데이터를 사용하는 Principled Technologies.^{19,20}

앞서 살펴본 것처럼 Supermicro 및 Dell GPU 기반 서버 오퍼링의 MLPerf® 성능 결과는 전체적으로 비슷하며, 한 AI 모델에서만 Dell PowerEdge XE8640 및 PowerEdge XE9640 서버가 유의미한 장점을 보였다.

AI를 구현하는 데 있어서 성능은 한 가지 요소일 뿐이므로 워크스테이션, 스토리지, 네트워킹부터 서비스, 지원, 교육 등까지 Dell과 Supermicro에서 제공하는 AI 오퍼링의 나머지 요소들도 살펴보았다. 그 결과, Dell AI 포트폴리오가 Supermicro 오퍼링보다 광범위하며 컴퓨팅 성능 외 여러 장애물에 대한 솔루션도 제공한다는 점을 확인하였다.

클라이언트 워크스테이션 및 스토리지 오퍼링

워크스테이션을 사용하는 추가 컴퓨팅 옵션

일부 AI 활용 사례에는 다른 컴퓨팅 접근 방식이 필요하며, AI 지원 하드웨어에 액세스해야 하는 모든 사람이 데이터 센터에 계속 연결되어 있을 수 있는 것은 아니다. 실험실에서 일하는 과학자들은 서버 랙을 설치할 공간이 없을 수 있으며, 오피스에서 일하는 사람들은 크고 무거운 데스크탑 시스템을 들고 다니는 것이 사실상 불가능하다. 그래서 워크스테이션이 필요한 것이다. Dell AI 포트폴리오에는 다양한 요구 사항을 충족하기 위해 타워, 모바일 및 랩 구성을 비롯한 여러 AI 지원 Precision 워크스테이션이 포함되어 있다.²¹ 반면에 Supermicro는 이동 중에도 쉽게 액세스할 수 있는 모바일 워크스테이션을 제공하지 않는다. GPU 워크스테이션 오퍼링은 여러 가지 타워 구성으로 구성되며, Supermicro는 이 중 일부를 사용자가 랙에 마운트할 수 있다고 주장한다.²² 연구 결과에 따르면 2개의 랙 마운트형 워크스테이션이 있지만, 이 워크스테이션들은 NVIDIA A100 GPU만 지원하는 구형 제품이며 아마도 더 이상 판매되고 있지 않는 것 같다.²³ 배포하는 GPU 워크스테이션의 유형과 이동성에 유연성이 필요한 기업이라면 Dell AI 포트폴리오가 더 적합하다.

스토리지 고려 사항

AI 워크로드를 실행하려면 컴퓨팅만큼이나 스토리지도 중요할 수 있다. 데이터가 많을수록 AI 모델 정확도가 향상되지만, 방대한 데이터 세트를 저장하고 관리하는 것은 많은 데이터 센터에서 부담이 될 수 있다. 또한 모델은 일반적으로 비정형 데이터를 사용하여 학습하기 때문에 AI 지원 스토리지 시스템은 다양한 데이터 유형을 쉽게 처리해야 한다.²⁴ 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.

Dell PowerScale 올플래시 NAS 포트폴리오는 노드당 3.84TB부터 720TB까지의 물리적 용량 옵션을 제공하며, 클러스터링된 올플래시 용량은 186PB의 물리적 용량에 달한다. PowerScale의 유연성과 확장성은 다양한 고객과 AI 활용 사례를 지원할 수 있다.²⁵ 세 가지 올플래시 PowerScale 모델(F200, F600, F900) 모두 스토리지 효율성을 높이기 위해 인라인 데이터 압축 및 중복 제거를 포함하고 있다.²⁶ 각 PowerScale 스토리지 모델은 Dell OneFS™ 파일 시스템을 사용하는데, 이 파일 시스템은 가장 자주 액세스하는 데이터가 가장 뛰어난 성능의 스토리지 계층에 상주하도록 고급 계층화 정책을 활용한다.²⁷ Dell은 APEX File Storage for AWS를 통해 AWS(Amazon Web Services) 마켓플레이스에서 OneFS 소프트웨어도 제공한다. 고객은 OneFS를 AWS 컴퓨팅 인스턴스와 함께 활용하여 온프레미스 OneFS 어레이에서와 동일한 기능을 사용할 수 있는 일관된 사용자 경험을 제공할 수 있다.²⁸

Supermicro 스토리지 오퍼링은 다양한 크기와 밀도의 스토리지 서버(랙 마운트형 고집적 스토리지 서버)로 구성된다.²⁹ Supermicro에서 파일 스토리지를 얻으려는 고객은 WekaIO, Scality Ring, OSNEXUS 등 다양한 타사 소프트웨어 정의 스토리지 오퍼링 중에서 선택해야 한다.³⁰ Scality Ring 및 OSNEXUS가 플랫폼 설명의 일부분으로 파일 스토리지 옵션을 포함하고는 있지만, 기본 파일 스토리지를 원하는 고객들은 주로 WekaIO를 선택하는 것으로 보인다. Supermicro는 다양한 활용 사례를 포괄하는 여러 레퍼런스 아키텍처를 제공하지만, 고객이 WekaIO 소프트웨어 구독 또는 라이선스를 가지고 있어야 하므로 전체적인 솔루션 비용은 증가할 수 있다.³¹

Dell의 오브젝트 스토리지 옵션에는 "퍼블릭 클라우드 규모로 비정형 데이터를 저장하도록 특별히 설계된" Dell ECS Enterprise Object Storage가 포함되어 있다.³² ECS 스토리지 노드는 하이브리드 클라우드 기능을 위한 Amazon S3 오브젝트 스토리지와의 기본 호환성과 함께 랙당 최대 14PB의 용량을 제공한다.³³ 파일 스토리지와 마찬가지로 Supermicro 오브젝트 스토리지 오퍼링에는 타사 구성이 필요하다. OSNEXUS 플랫폼은 파일, 블록 및 오브젝트 스토리지를 위한 통합 스토리지 플랫폼인 반면 Scality RING 솔루션은 파일 및 오브젝트 스토리지를 결합한다. 이 두 제품 모두 타사 공급업체와 라이선스를 체결해야 한다.^{34,35} 오브젝트 스토리지의 경우에 한해서만 고객은 Quantum 소프트웨어 구독을 통해 프라이빗 클라우드 오브젝트 스토리지를 위한 Quantum ActiveScale용 Supermicro 솔루션을 구매할 수 있다.³⁶ 이 문서를 작성하는 시점에서는 Supermicro가 제공하는 유연한 소비/운영 지출 옵션을 찾을 수 없었다.

Supermicro 스토리지 오퍼링은 고객이 타사 소프트웨어 공급업체와 계약을 체결해야 하기 때문에 추가 라이선싱 또는 구독 비용이 발생할 수 있으며 지원, 문제 해결 등에 어려움을 겪을 수 있다. Dell AI 포트폴리오는 Dell 스토리지 고객에게 스토리지 솔루션의 모든 측면에 걸쳐 신뢰할 수 있는 단일 서비스 및 지원 솔루션을 제공한다.

Dell APEX 정보

파일 스토리지를 서비스로 이용하려는 고객을 위해 Dell은 파일, 블록 및 백업 스토리지를 포함하는 APEX Data Storage Services를 제공한다. 고객은 Dell APEX Console을 사용하여 새 구독을 주문하고, 스토리지 용량을 조정하고 모니터링하는 등의 작업을 수행할 수 있다. Dell에 따르면 이 솔루션을 사용함으로써 "애플리케이션과 데이터를 더 효과적으로 제어하면서 클라우드 경험의 편의성과 민첩성을 얻을 수 있다"고 한다.³⁷

Dell APEX에 대한 자세한 내용은 <https://www.dell.com/en-us/dt/apex/storage/data-storage-services/index.htm>을 참조하면 된다.

네트워킹 옵션

네트워킹은 AI 인프라스트럭처의 또 다른 중요한 구성 요소이다. 많은 AI 워크로드는 서로 간에 그리고 스토리지와의 지속적인 통신이 필요한 대규모 서버 클러스터에서 실행되므로, 병목 현상을 방지하기 위해 AI 워크로드에 강력한 네트워킹이 필요하다. AI 워크로드에 대한 네트워킹이 부족하면 학습 및 추론 시간이 늘어나서 데이터 처리 속도가 느려지고 통찰력을 확보하기까지 오래 걸린다. Dell은 이더넷 및 패브릭 네트워킹을 위한 PowerSwitch Data Center ToR(Top-of-Rack) 스위치와 PowerEdge MX I/O 모듈을 제공한다.³⁸ PowerSwitch는 다양한 요구에 부합하기 위해 1GbE부터 400GbE까지 다양한 오퍼링을 제공한다. 또한 Dell PowerSwitch Z Series 스위치는 리프/스파인 패브릭에 최적화된 100GbE 및 400GbE 연결을 제공한다.³⁹

Supermicro는 ToR 그리고 데이터 센터 스파인 및 리프 등의 기타 애플리케이션을 위해 최대 400GbE의 이더넷 포트를 갖춘 스위치도 제공한다.⁴⁰ 그러나 Dell 네트워킹 서비스는 Supermicro가 제공하지 않는 여러 가지 사용 편의성과 유연성 이점을 제공한다. Dell Fabric Design Center와 같은 서비스는 네트워킹 불일치, 격차 또는 비효율성을 방지하여 고객이 자동화를 통해 네트워크 패브릭을 계획하고 배포할 수 있도록 도와준다.⁴¹ VMware VxRail, VMware ESXi 및 Dell PowerStore 구성과 같은 특정 환경을 위해 Dell은 소프트웨어 정의 인프라스트럭처 배포 및 수명주기 관리를 지원하는 SmartFabric Services를 제공한다. SmartFabric Services는 PowerStore를 통해 "플러그 앤 플레이 패브릭으로 LAN 연결 작업의 최대 99%"를 자동화할 수 있다.⁴² 네트워킹 설계 및 구현을 위해 자동화, 지침 등을 제공하는 이 서비스는 고객이 AI를 도입하는 과정을 지원한다.

서비스, 교육 등

AI를 구현할 때 하드웨어를 제외한 주요 과제는 전략, 계획, 데이터 준비 및 관리를 위한 사내 전문 지식이 확보해야 한다는 것이다. AI 워크로드를 관리하고 유지하려면 기존의 하드웨어 전문 지식과 머신러닝 운영 및 데이터 사이언스를 모두 포함하는 특별한 지식이 필요하다. 또한 AI 전략을 설계하고 구현하는 사람들은 새로운 AI 워크로드가 이러한 목표에 부합하도록 하기 위해 회사의 고유한 운영 목표를 심층적으로 이해하고 있어야 한다.⁴³

또 다른 커다란 장애물은 AI를 기존 운영 시스템에 원활하게 통합하는 것이다. 이러한 통합을 위해서는 새로운 AI 기술을 현재 비즈니스 프로세스와 전략적으로 연계함으로써 AI를 도입해도 기존 워크플로는 그대로 유지되도록 해야 한다. 최적화되고 검증된 여러 솔루션 레퍼런스 아키텍처, 교육 과정, 관리 옵션 및 대규모 파트너 생태계를 제공하는 Dell과 같은 회사와 파트너십을 맺으면 AI를 도입하는 것이 쉬워질 수 있다.

AI를 위한 전문 서비스

교육 및 계획을 위해 Dell은 다양한 AI 전용 서비스를 제공한다.⁴⁴ AI 구현을 지원하는 Dell 서비스에는 컨설팅, 데이터 준비, 배포, 지원 및 교육이 포함되며, 각 서비스는 AI 도입의 특정 측면을 대상으로 한다. Dell Advisory Services for Generative AI는 고객이 활용 사례를 식별하고 기업 프로세스를 간소화하는 데 도움이 되는 로드맵을 만들 수 있도록 지원한다.⁴⁵ 마찬가지로 Adoption Services for Generative AI는 Dell 전문가와의 워크숍을 통해 고객의 요구 사항과 고유한 당면 과제를 검토하여 비즈니스에 맞는 사전 학습된 모델을 결정하고 IT 직원을 교육하는 지식 전달 세션을 실시한다.⁴⁶ 또한 Dell의 Implementation, Scale, and Managed Services for Generative AI는 완전히 관리되는 AI 인프라스트럭처에 이르기까지 다양한 수준의 지원과 교육을 제공한다. 따라서 고객사의 IT 직원은 하드웨어 관리를 Dell에 맡기고 모델과 데이터에 집중할 수 있다.⁴⁷ ProSupport Services는 최적의 시스템 성능을 보장하며, 지속적인 AI 운영을 위한 필수 하드웨어 및 소프트웨어 지원을 제공하여 기술적 문제를 해결한다.⁴⁸

교육 서비스는 AI 활용에 필요한 기술과 지식을 육성하는 데 필수적이다. Dell 교육 오퍼링에는 데이터 사이언스에 대한 포괄적인 교육 프로그램, 고급 분석 인증 그리고 머신 러닝과 같은 특정 AI 기술에 대한 워크숍이 포함된다.⁴⁹

반면에 Supermicro 서비스는 대부분 문제 해결, 설명서, RMA(Return Merchandise Authorization) 및 보증으로 제한되어 있다.⁵⁰ Supermicro AI 포트폴리오에서는 설계, 구현, 관리 또는 교육 서비스를 찾을 수 없었다. 복잡한 AI 도입 과정을 헤쳐 나가기 위한 교육 파트너를 찾는 기업들에게는 이 두 회사 중 Dell이 명확한 선택지이다.

AI 워크로드를 위한 타사 파트너십

Dell Technologies와 NVIDIA는 비즈니스 환경에서 Generative AI를 위한 포괄적인 솔루션을 제공하는 것을 목표로 하는 Dell Validated Designs를 공동으로 제공한다. 이 프로젝트는 기업이 맞춤형 AI 모델을 구축하고 운영할 수 있도록 지원하는 AI 모델 프레임워크와 더불어 Dell 및 NVIDIA 기술 및 소프트웨어를 기반으로 하는 확장형 고성능 인프라스트럭처를 구축한다. 이 솔루션을 통해 고객은 GenAI 워크로드를 신속하게 가동할 수 있다.⁵¹ 자세한 내용은 아래의 Dell Validated Designs 섹션을 참조하면 된다.

Dell은 AI 기술 애플리케이션을 개선하기 위해 여러 회사와 협력하고 있다. Hugging Face와 함께 Dell은 사이트에서 LLM(Large Language Model)을 설정하는 것을 더 쉽게 만들고 있다. 이 파트너십은 Hugging Face AI 전문 지식과 Dell 서버 및 스토리지 시스템을 결합한다. Dell 전용 Hugging Face 포털은 Hugging Face 오픈 소스 AI 모델을 간단하고 안전하게 배포할 수 있는 툴을 제공할 것이다. 지속적인 목표는 Dell 시스템을 위해 이러한 모델을 계속 개선하면서 성능을 높이고 새로운 AI 애플리케이션을 지원하는 것이다.⁵²

Dell과 Starburst는 Starburst 분석을 Dell 컴퓨팅 및 스토리지 기술과 통합하는 확장형 고성능 데이터 레이크를 개발하고 있으며, AI 및 ML 툴의 모든 데이터 소스에 대한 단일 액세스 지점을 제공하고자 한다. 고객은 이 파트너십을 활용하여 데이터 사일로 제거할 수 있게 될 것이다.⁵³

연구에 따르면 Supermicro는 AI를 위한 파트너십이 훨씬 제한적이다. SiMa.ai와 Supermicro는 멀티 스트림 비디오 분석 처리를 위해 설계된 컴팩트 엣지 ML 서버인 Supermicro SYS-E300-13AD를 개발하기 위해 협력했다. 칩에 SiMa.ai ML 파이프라인을 탑재한 이 서버는 여러 비디오 채널을 효율적으로 처리하고, 총 소유 비용을 줄이고, 신뢰성과 보안을 강화한다. 이 서버는 수많은 비디오 스트림을 처리하고 분석하도록 설계된 컴퓨팅 설정을 통해 다양한 엔터프라이즈 애플리케이션에 적합한 엣지 인텔리전스를 제공한다.⁵⁴

Dell Validated Designs

AI 하드웨어 솔루션에서 불확실성을 해소할 수 있도록 Dell은 랩 검증을 거쳐 여러 AI 및 기타 워크로드에 최적화된 것으로 입증된 레퍼런스 아키텍처를 제공한다. 이러한 Validated Designs에는 아키텍처 개념, 전체 솔루션 개요 및 성능 그리고 설계 시 고려한 워크로드에 대해 솔루션의 기능을 입증하는 기타 랩 검증이 포함되어 있다. 이러한 워크로드에는 가상화 환경, MLOps, 머신 러닝, 대화형 AI, GenAI 추론, GenAI 모델 튜닝, NVIDIA Fleet Command, OpenShift AI가 포함된다.⁵⁵

예를 들어, AI for Virtualized Environments Validated Design은 Dell 인프라스트럭처에서 VMware 지원 AI와 NVIDIA AI Enterprise를 결합하여 가상 설정에서 AI를 최적화한다.⁵⁶ Validated Design 가이드에는 ResNet 모델 학습 성과를 보여주는 성능 결과가 포함되어 있다. 이 결과는 고객들에게 설계가 어떻게 작동하는지를 증명하고 어떤 성능을 기대할 수 있는지를 보여준다.⁵⁷ 이러한 검증은 개념을 설명하고, 구성을 추천하고, 고려 사항과 성능 기대치를 고객에게 안내함으로써 단지 어떤 하드웨어가 서로 연동되는지를 나열하는 것이 아니라 그 이상의 가치를 고객에게 제공한다.⁵⁸

Supermicro도 활용 사례를 기반으로 솔루션을 제공하지만 Dell Validated Designs에서 확인한 것과 같은 수준의 레퍼런스 아키텍처에는 미치지 못한다. Supermicro는 특정 워크로드를 위해 특별히 설계된 솔루션을 추천하는 대신 서버와 GPU를 AI 추론 및 학습, HPC/AI, 시각화, 설계 등의 범주로 정리하고 있다.⁵⁹ 브로셔 및 데이터시트에서 이러한 범주는 Supermicro가 작업에 가장 적합한 것으로 추천하는 몇몇 서버와 GPU, 여러 활용 사례, 주요 사용 기술 및 제안 소프트웨어 목록으로 구성되어 있다.⁶⁰ Dell Validated Designs와 달리 네트워크 아키텍처와 성능 또는 검증 데이터는 포함되어 있지 않는 것으로 보인다. 또한 Supermicro는 NVIDIA와 함께 게시한 액체 냉각 솔루션 브리프를 사용한 대규모 AI 학습 등 일부 AI 솔루션에 대해 보다 상세한 레퍼런스 아키텍처를 제공하는 여러 Reference Designs를 제공한다.⁶¹ 고객은 특정 시나리오에 대해 보다 심층적인 레퍼런스 아키텍처를 찾을 수 있지만, 연구 시점 당시에는 앞서 언급한 액체 냉각 아키텍처, AI 워크스테이션 아키텍처, RedHat OpenShift 아키텍처의 세 가지만 찾을 수 있었다.⁶²

전반적으로 Dell Validated Designs가 Supermicro 오퍼링보다 더 많은 AI 워크로드를 다루고 더 심층적인 지침을 제공한다는 것이 확인되었다.



관리 서비스 및 iDRAC

Principled Technologies의 **2023년 4월 보고서**에 따르면 iDRAC(integrated Dell Remote Access Controller)은 특히 자동화, 보안 및 구성에서 Supermicro IPMI(Intelligent Platform Management Interface)에 비해 여러 가지 고급 기능을 제공한다.⁶³ 표 3은 iDRAC이 Supermicro IPMI보다 배포 및 펌웨어 업데이트가 더 쉽고 더 많은 보안 기능을 제공하는 것을 보여주는 보고서에서 발췌한 Dell 및 Supermicro 관리 기능의 비교이다. 일부 결과는 최초 게시 이후 바뀌었을 수 있다.

표 3: Dell 관리 툴과 Supermicro 관리 툴을 비교한 2023년 4월 Principled Technologies 결과의 요약. 일부 결과는 게시 이후 바뀌었을 수 있다.
출처: Principled Technologies <https://facts.pt/V5fDf06>.

	Dell 관리 툴의 차이점	우수성
더 간편한 펌웨어 업데이트 <i>iDRAC9과 Supermicro IPMI 비교</i> <i>OME와 Supermicro SSM 비교</i>	- 예약 옵션을 통한 iDRAC9 자동 온라인 업데이트 - OME를 사용하면 맞춤형 펌웨어 리포지토리를 생성할 수 있으며 추가 툴이나 에이전트 없이 BIOS, BMC 및 기타 서버 구성 요소의 펌웨어를 업데이트할 수 있음	- 단 74초 만에 iDRAC에서 자동 업데이트를 설정했음 - Supermicro IPMI에는 자동 업데이트 기능이 없으므로 관리자가 수동으로 업데이트해야 함 - SSM은 BIOS 및 BMC 펌웨어 업데이트만 지원하며 다른 구성 요소를 업데이트하려면 SUM이 필요함
더 많은 보안 기능 <i>iDRAC9과 Supermicro IPMI 비교</i> <i>OME와 Supermicro SSM 비교</i>	- iDRAC9은 MFA를 제공하며 시스템 다운타임 없이 USB 포트를 동적으로 비활성화할 수 있음 - OME는 RBAC(Role-Based Access Control)와 SBAC(Scope-Based Access Control)을 모두 제공하여 디바이스 관리를 디바이스 그룹의 하위 집합으로 제한함	- Supermicro IPMI에는 MFA 기능이 없음 - Supermicro IPMI를 사용하려면 시스템을 재부팅하고 BIOS 구성으로 들어가서 USB 포트를 비활성화해야 함 - Supermicro SSM은 RBAC를 제공하지만 더 제한적인 SBAC는 제공하지 않음
더 간편한 수명주기 관리 <i>OME와 Supermicro SSM 비교</i>	OME를 통한 에이전트 없는 전체 수명주기 관리로 관리 및 모니터링 간소화	SSM을 사용하려면 자세한 로컬 시스템 상태 메트릭을 위한 SuperDoctor5 에이전트와 추가 구성 요소를 업데이트하기 위한 SUM(Supermicro Update Manager)이 필요함
더 간편한 서버 구축 <i>iDRAC9과 Supermicro IPMI 비교</i>	- iDRAC9을 사용하여 단 12 단계 만에 전체 Dell 서버 프로파일을 가져올 수 있음 52가지 BIOS 기능을 제공하는 iDRAC9을 통한 강력한 BIOS 구성 옵션 그리고 RAID NIC 및 iDRAC과 같은 구성 요소 구성에 대한 지원	- Supermicro IPMI를 사용하면 전체 서버 프로파일이 아닌 IPMI 구성만 저장하고 복원할 수 있음 - iDRAC9에는 52개의 BIOS 기능이 있지만 IPMI는 BIOS 구성 옵션 제공하지 않음
더 많은 보고 및 분석 옵션 <i>iDRAC9과 Supermicro IPMI 비교</i> <i>OME와 Supermicro SSM 비교</i>	- iDRAC9은 텔레메트리 스트리밍을 제공하여 사용자가 서버 데이터를 Splunk와 같은 분석 툴에 쉽게 전송할 수 있도록 함 - OME는 좀 더 쉬운 모니터링을 위해 텔레메트리 데이터를 CloudIQ로 직접 전송함	- IPMI는 관리자가 집계 및 최종 분석을 위해 메시지를 보내는 데 사용할 수 있는 syslog 기능만 제공함 - SSM에는 Dell CloudIQ와 동등한 클라우드 기반 관리 솔루션이 없음
더 많은 지속 가능성 기능 <i>OME와 Supermicro SSM 비교</i>	OME Power Manager에는 모니터링을 위한 더 많은 메트릭이 있음(탄소 배출량 데이터 포함)	SSM은 활용도 메트릭이 더 적으며 탄소 배출량을 추적할 방법이 없음
더 많은 모니터링 방법 <i>OME와 Supermicro SSM 비교</i>	- OpenManage Mobile 앱을 통해 어디서든 Dell 서버 관리 - 타사 SNMP MIBS 가져오기를 지원하는 서버 IP 및 자격 증명을 사용하여 OME로 타사 디바이스 모니터링	- SSM에는 모바일 앱이 없음 - SSM은 서버 IP를 사용하여 타사 디바이스를 모니터링하는 것을 허용하지 않음



결론

비즈니스에서 AI 솔루션을 설계, 구현, 관리 및 유지 보수할 때 고려해야 할 여러 가지 요소들이 있다. 현명하게 투자하고 AI 솔루션을 최대한 활용하려면 하드웨어 공급업체 이상의 역할을 할 수 있는 공급업체를 찾아야 한다. 연구 결과, Dell이 이 과정 전반에서 파트너로서 도움이 될 수 있는 서비스를 제공하는 것으로 나타났다. 따라서 AI에 대한 투자 계획을 세울 때는 Dell과 함께하는 것을 고려해 보는 것이 좋다.

1. Vipera, "NVIDIA's H100 and A100 GPU Cards: Exploring the Intricacies of SXM and PCI-E Connections", 2024년 1월 5일 열람, <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", 2024년 1월 5일 열람 <https://mlcommons.org/en/inference-datacenter-31/>.
3. MLCommons, "MLPerf Inference: Datacenter Benchmark Suite Results".
4. Dell, "PowerEdge XE 서버", 2024년 1월 5일 열람, <https://www.dell.com/en-us/dt/servers/specialty-servers/poweredge-xe-servers.htm>.
5. Supermicro, "Next Leap of AI Infrastructure is Here", 2024년 1월 5일 열람, <https://www.supermicro.com/en/accelerators/nvidia>.
6. Dell, "PowerEdge XE9680", 2024년 1월 5일 열람, <https://www.delltechnologies.com/asset/en-us/products/servers/technical-support/poweredge-xe9680-spec-sheet.pdf>.
7. 검증된 MLPerf® v3.1 Inference Closed 점수. <https://mlcommons.org/benchmarks/inference-datacenter/> 2023년 12월 5일, 3.1-0069 항목에서 발체. 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.
8. 검증된 MLPerf® v3.1 Inference Closed 점수. <https://mlcommons.org/benchmarks/inference-datacenter/> 2023년 12월 5일, 3.1-0135 항목에서 발체. 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.
9. 검증된 MLPerf® v3.1 Inference Closed 점수. <https://mlcommons.org/benchmarks/inference-datacenter/> 2023년 12월 5일, 3.1-0069 항목에서 발체. 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.
10. 검증된 MLPerf® v3.1 Inference Closed 점수. <https://mlcommons.org/benchmarks/inference-datacenter/> 2023년 12월 5일, 3.1-0132 항목에서 발체. 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.
11. Dell, "PowerEdge XE9640 랙 서버", 2024년 1월 5일 열람, <https://www.dell.com/en-us/shop/ipovw/poweredge-xe9640>.
12. Accelsius, "Enabling the AI Revolution with Liquid Cooling", 2024년 1월 5일 열람, <https://www.accelsius.com/blog/enabling-the-ai-revolution-with-liquid-cooling>.

13. Dell, "PowerEdge XE9640 기술 가이드", 2024년 1월 5일 열람, <https://www.delltechnologies.com/asset/en-us/products/servers/technical-support/powerededge-xe9640-technical-guide.pdf>.
14. Supermicro, "GPU Server Systems", 2024년 1월 5일 열람, https://www.supermicro.com/en/products/gpu?pro=pl_grp_type%3D1.
15. 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.org를 참조하십시오.
16. 검증된 MLPerf® v3.1 Inference Closed 점수. <https://mlcommons.org/benchmarks/inference-datacenter/> 2023년 12월 5일, 3.1-0133 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.org를 참조하십시오.
17. Dell, "PowerEdge XE8640", 2024년 1월 5일 열람, <https://www.delltechnologies.com/asset/en-us/products/servers/technical-support/powerededge-xe8640-spec-sheet.pdf>.
18. Dell, "PowerEdge XE8640 랙 서버", 2024년 1월 5일 열람, <https://www.dell.com/en-us/shop/ipovw/powerededge-xe8640>.
19. 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.org를 참조하십시오.
20. 검증된 MLPerf® v3.1 Inference Closed 점수. <https://mlcommons.org/benchmarks/inference-datacenter/> 2023년 12월 5일, 3.1-0133 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.org를 참조하십시오.
21. Dell, "Dell Precision 워크스테이션 기반 AI(Artificial Intelligence) 기술", 2024년 1월 5일 열람, <https://www.dell.com/en-us/dt/ai-technologies/index.htm?hve=explore+dell+precision+for+ai#tab0=0>.
22. Supermicro, "Super Workstations", 2024년 1월 5일 열람, <https://www.supermicro.com/en/products/superworkstation>.
23. Supermicro, "Rackmount Workstations", 2024년 1월 5일 열람, <https://www.supermicro.com/en/products/rackmount-workstations>.
24. Stephen Pritchard, "Storage requirements for AI, ML and analytics in 2022", 2024년 1월 5일 열람, <https://www.computerweekly.com/feature/Storage-requirements-for-AI-ML-and-analytics-in-2022>.
25. Dell, "PowerScale AI 지원 데이터 플랫폼", 2024년 1월 5일 열람, <https://www.dell.com/en-us/shop/powerscale-family/sf/powerscale>.
26. Dell, "Dell PowerScale All-Flash"(2024년 01월 05일 액세스), <https://www.delltechnologies.com/asset/en-us/products/storage/technical-support/h15963-ss-powerscale-all-flash-nodes.pdf>.
27. Dell, "Dell PowerScale OneFS 소프트웨어 기능"(2024년 01월 05일 액세스), <https://www.delltechnologies.com/asset/en-us/products/storage/technical-support/h18275-onefs-software-features-data-sheet.pdf>.
28. Dell, "Dell ECS 엔터프라이즈 오브젝트 스토리지", 2024년 1월 5일 열람, <https://www.dell.com/en-us/dt/storage/ecs/>.
29. Supermicro, "Accelerating AI Data Pipelines", 2024년 1월 5일 열람, <https://www.supermicro.com/en/products/storage>.
30. Supermicro, "Supermicro Software-Defined Storage and Memory Solutions", 2024년 1월 5일 열람, <https://www.supermicro.com/en/solutions/software-defined-storage>.
31. Supermicro, "Supermicro WEKA Distributed Storage Solution", 2024년 1월 5일 열람, <https://www.supermicro.com/en/solutions/wekaio>.
32. Dell, "Dell ECS 엔터프라이즈 오브젝트 스토리지", 2024년 1월 5일 열람, <https://www.dell.com/en-us/dt/storage/ecs/>.
33. Dell, "Dell ECS Enterprise Object Storage".
34. Supermicro, "Supermicro OSNEXUS Software-Defined Storage Solution", 2024년 1월 5일 열람, <https://www.supermicro.com/en/solutions/osnexus>.
35. ASBIS, "Supermicro Solution for Scalify RING", 2024년 1월 5일 열람, <https://news.asbis.com/news/suppliers/supermicro-renewed-the-line-of-scality-ring-solution/>.
36. Supermicro, "Supermicro solution for Quantum ActiveScale", 2024년 1월 5일 열람, <https://www.supermicro.com/en/solutions/activescale>.
37. Dell, "확장 가능하고 탄력적인 Storage as-a-Service", 2024년 1월 5일 열람, <https://www.dell.com/en-us/dt/apex/storage/data-storage-services/>.
38. Dell, "PowerSwitch를 통해 오픈 네트워킹으로 전환", 2024년 1월 5일 열람, <https://www.dell.com/en-us/dt/networking/>.

-
39. Dell, "Dell PowerSwitch 데이터 센터 스위치", 2024년 1월 5일 열람, <https://www.dell.com/en-us/dt/networking/data-center-switches/>.
 40. Supermicro, "SSE-T7132S - 400Gb Ethernet Switch", 2024년 1월 5일 열람, <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", 2024년 1월 5일 열람, <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", 2024년 1월 5일 열람, <https://www.dell.com/en-us/dt/networking/smartfabric/>.
 43. Penny Madsen, "Scaling Skills for AI: Lessons from Early Adopters", 2024년 1월 5일 열람, <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", 2024년 1월 5일 열람, <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", 2024년 1월 5일 열람, <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", 2024년 1월 5일 열람, <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", 2024년 1월 5일 열람, <https://www.dell.com/en-us/dt/services/solutions/artificial-intelligence-services.htm>.
 49. Dell, "Comprehensive AI Training Modules Tailored for You", 2024년 1월 5일 열람, <https://education.dell.com/content/emc/en-us/home/training/aiml.html>.
 50. Supermicro, "Services and Support", 2024년 1월 5일 열람, <https://www.supermicro.com/en/support>.
 51. Travis Vigil, "Dell and NVIDIA: Bringing Generative AI to the Enterprise", 2024년 1월 5일 열람, <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", 2024년 1월 5일 열람, <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", 2024년 1월 5일 열람, <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", 2024년 1월 5일 열람, <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 솔루션", 2024년 1월 5일 열람, <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", 2024년 1월 5일 열람, <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", 2024년 1월 5일 열람, <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", 2024년 1월 5일 열람, <https://infohub.delltechnologies.com/l/design-guide-virtualizing-gpus-for-ai-with-vmware-and-nvidia-based-on-dell-infrastructure-1/design-considerations-105/>.
 59. Supermicro, "Accelerate Every Workload", 2024년 1월 5일 열람, <https://www.supermicro.com/en/solutions/ai-deep-learning>.
 60. Supermicro, "Supermicro Enterprise AI Inference & Training", 2024년 1월 5일 열람, 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", 2024년 1월 5일 열람, https://www.supermicro.com/solutions/Solution-Brief_Rack_Scale_AI.pdf.
62. Supermicro, "Accelerate Every Workload", 2024년 1월 5일 열람, <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", 2024년 1월 5일 열람, <https://www.principledtechnologies.com/Dell/Management-tools-vs-Supermicro-0423.pdf>.

▶ 이 보고서의 원본(영문) 보기:
<https://facts.pt/q9p46K9>

이 프로젝트는 Dell Technologies의 의뢰로 진행되었습니다.



Facts matter.®

Principled Technologies는 Principled Technologies, Inc.의 등록 상표입니다.
기타 모든 제품 이름은 해당 소유자의 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 명시적 보증을 비롯하여 테스트 결과 및 분석, 정확성, 완전성 또는 품질에 대한 어떠한 명시적 또는 묵시적 보증에 대해서도 명확하게 부인합니다. 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 손해, 특수한 손해, 부수적 손해 또는 결과적 손해에 대해 책임지지 않으며, 이는 사전에 그러한 손해의 가능성을 통보받은 경우에도 마찬가지입니다. 어떠한 경우에도 직접적 손해를 포함한 Principled Technologies, Inc.의 책임은 Principled Technologies, Inc.에 테스트와 관련하여 지불된 비용을 초과하지 않습니다. 고객에 대한 유일하고 배타적인 구제책은 여기에 명시된 내용을 따릅니다.