



TEN

회사소개서

Make AI accessible for all.

Make AI accessible for all.



회사명	주식회사 텐 TEN Inc.
대표이사	오세진
설립일	2020.07.03
주소	서울특별시 강남구 테헤란로27길 18, 5층
주요 서비스	MLOps 솔루션 'AI Pub' AI 인프라 컨설팅 'RA:X'
직원 수	27명
홈페이지 / SNS	https://ten1010.io/ https://www.youtube.com/@ten1300 https://www.linkedin.com/company/ten1010/

2020년 서비스를 시작하여 각종 사업 선정과 어워드 수상으로 경쟁력과 잠재력을 인정받고 있습니다.

2020

- 02 AI Pub 서비스로 비즈니스 시작
- 07 주식회사 텐 법인 설립

2021

- 03 NVIDIA Inception Program 체결
- 04 Kingsley Ventures Seed 투자 유치
- 06 벤처기업 인증 획득
- 07 TIPS 창업성장기술개발사업 선정
- 08 기업 부설 연구소 설립
- 08 대한민국 리딩기업대상 K-스타트업대상 'AI 서비스 부문' 수상
- 11 NVIDIA GTC³ Top Startup from Korea 선정

2022

- 03 AI 바우처 사업 선정
- 07 대한민국 리딩기업대상 K-스타트업대상 'MLOps 플랫폼 부문'
- 11 한국평가데이터 기술역량 우수기업 인증 (T4 등급)
- 12 Pre-series A 라운드 투자(30억 규모) 유치

2023

- 01 가상화 리소스 관리, 크기 추천 방법 등 특허 5개 등록
- 02 TEN 제1회 AI 워크샵 개최
- 07 NetApp Preferred Partnership 체결
- 08 IBM Partnership 체결
- 09 Redhat Partnership 체결
- 11 한국국방과학연구소 딥러닝 연산용 고밀도 GPU기반 클러스터 컴퓨팅 시스템 사업 구축 완료
- 11 TEN 제2회 AI워크샵 개최
- 12 AI Pub GS 인증 획득
- 12 CIO Review '2023년 가장 유망한 한국 테크기업' 선정

2024

- 01 성균관대 SAL · COMPASS LAB AI 인프라 산학협력 체결
- 04 코오롱 베니트 AI Alliance MOU 체결
- 05 제19회 디지털 이노베이션 대상 IT부문 수상

2025

- 05 Series A 라운드 투자(80억 규모) 유치
- 05 2025 디지털이노베이션 대상 선정
- 08 NVIDIA Partner Network – Solution Advisor 체결



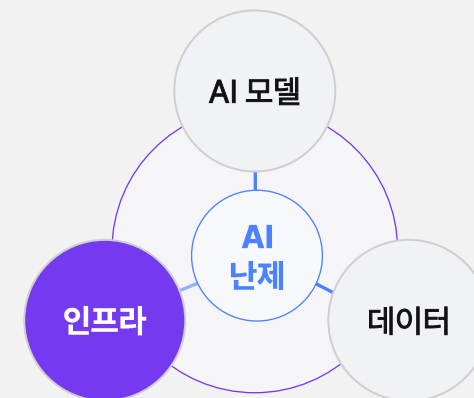
주식회사 텐은

Mission

세상을 널리 AI롭게 하기 위해

Vision 2024

기술의 장벽을 낮추는 소프트웨어 도구를 만듭니다.



특정 지식인이 아닌 모든 이가 AI를 만든다는 인력 보편화의 조건과
자신의 지식을 데이터로 구체화하면 AI를 만들 수 있다는 데이터 보편화의 조건은 점차 해결되고 있지만,
인프라는 철저하게 **비용**의 문제라 시간과 노력만으로는 해결하기 어렵습니다.

인력

머신러닝 박사 학위가 없어도
AI를 만들 수 있는 시대

데이터

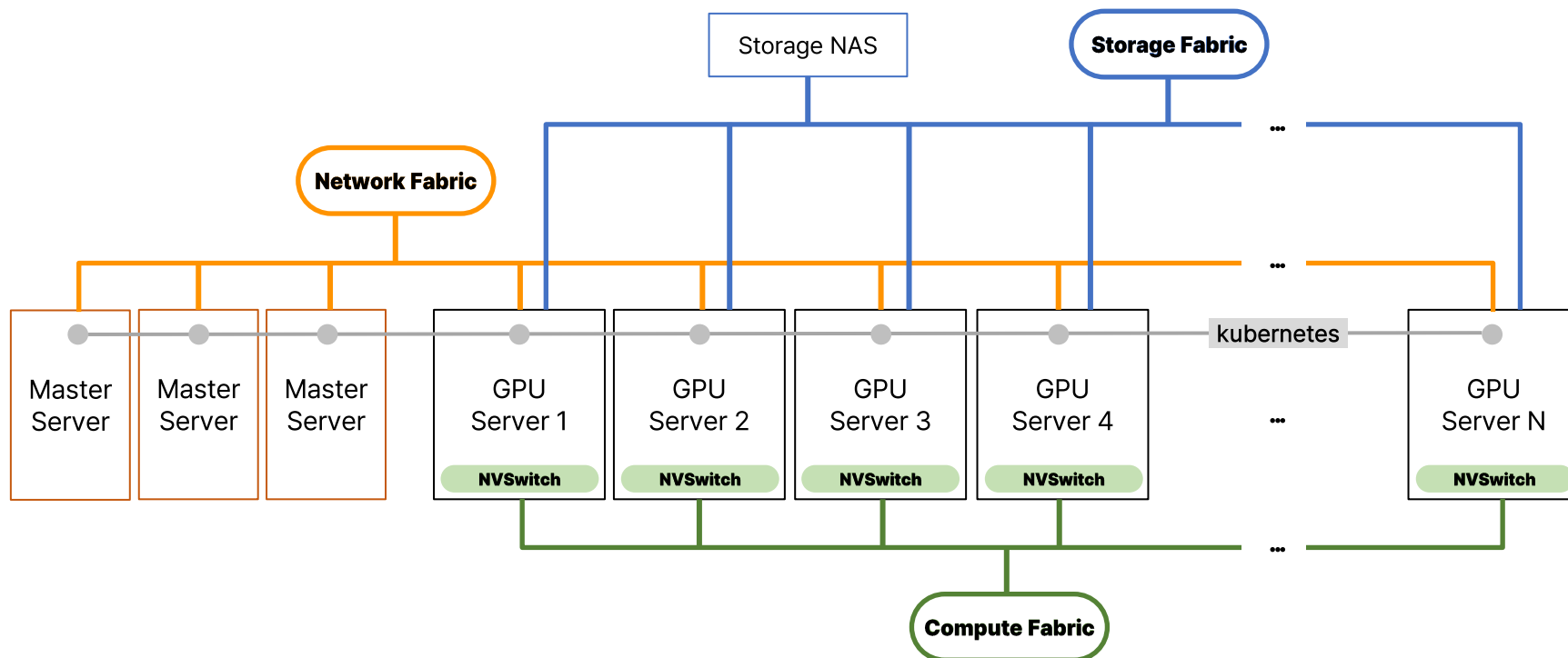
누구나 자신의 지식을
데이터로 구체화하고
이를 학습시킬 수 있음

인프라

인프라 = 비용

그러나 인프라 자원을 효율화하는 것은 어려운 일입니다.
AI 모델의 자원 수요가 지속적으로 증가함에 따라, 인프라 구성 역시 복잡해지기 때문입니다.

AI 모델 계산량 증가를 커버하기 위해, 다수의 GPU 머신을 도입할 수밖에 없고
이에 따른 AI 전용 인프라 구성은 더욱 복잡하고 어려워집니다.





**인프라 자원의 효율화는
AI 보편화의 핵심입니다.**

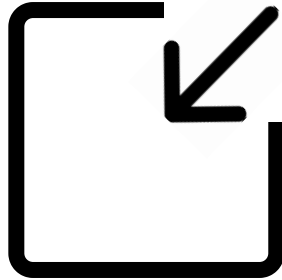


AI 구현은 얼마나 효율적으로 AI 인프라를 구축하고 쓸 수 있는가에 좌우됨

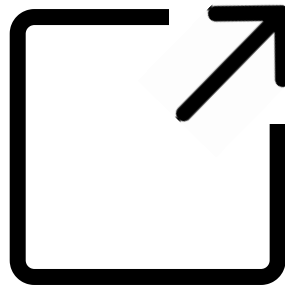
AI의 등장은 기존 컴퓨팅 패러다임의 혁신적 변화를 초래했음

혁신적 변화의 핵심은 GPU 뿐 아니라 다른 HW 요소들 포함한 인프라의 통합적 성능 향상이 필수적이라는 것임

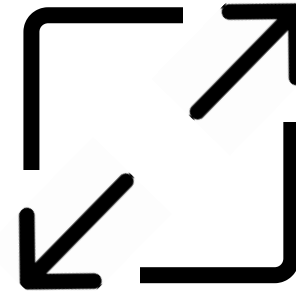
[AI로 인한 컴퓨팅 패러다임의 변혁]



서버 **내부**의 변혁



서버 **외부**의 변혁

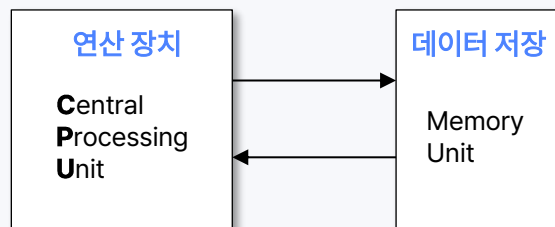


자원**분배**방식의 변혁

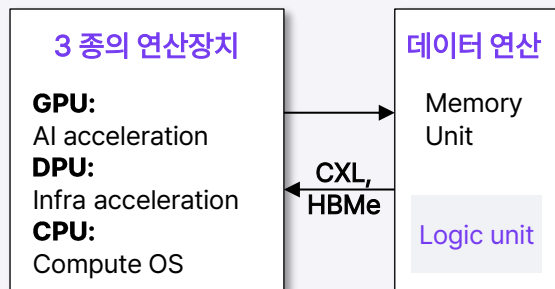
[AI로 인한 컴퓨팅 패러다임의 변혁 - 상세]

서버 내부의 변혁

폰 노이만 구조
과거 컴퓨터 시스템의 기본 구조

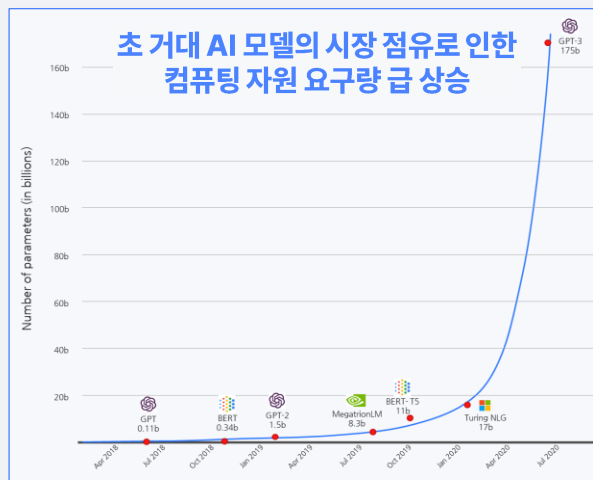


AI 데이터 센터 용 컴퓨터 시스템 구조
각기 다른 연산을 담당하는 3개의 처리 장치 & 연산용 프로세스를 더한 메모리 구조

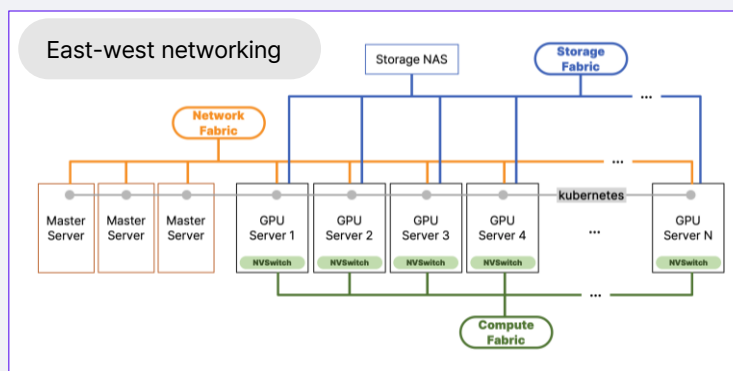


서버 외부의 변혁

**초 거대 AI 모델의 시장 점유로 인한
컴퓨팅 자원 요구량 급 상승**



초 거대 AI 모델의 계산량을 커버하기 위해 다수의 서버를
클러스터로 구성하며, 그 사용 방법이 복잡해지고 있음



자원 분배 방식의 변혁

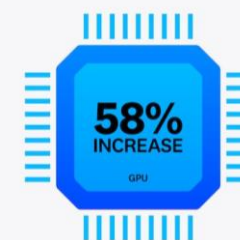
컨테이너 기술이 빠르게 확산함에 따라
컨테이너 플랫폼 쿠버네티스가
모던 인프라기술의 사실상 표준으로 부상



컨테이너화된 GPU 컴퓨팅이 YoY 58% 증가
AI의 방대한 데이터 처리 규모로 인해
컨테이너를 기반으로 한 GPU 컴퓨팅이 빠르게 성장 중

**CONTAINERIZED
GPU USAGE IS
ON THE RISE**

Year-over-year growth of containerized
GPU instance hours

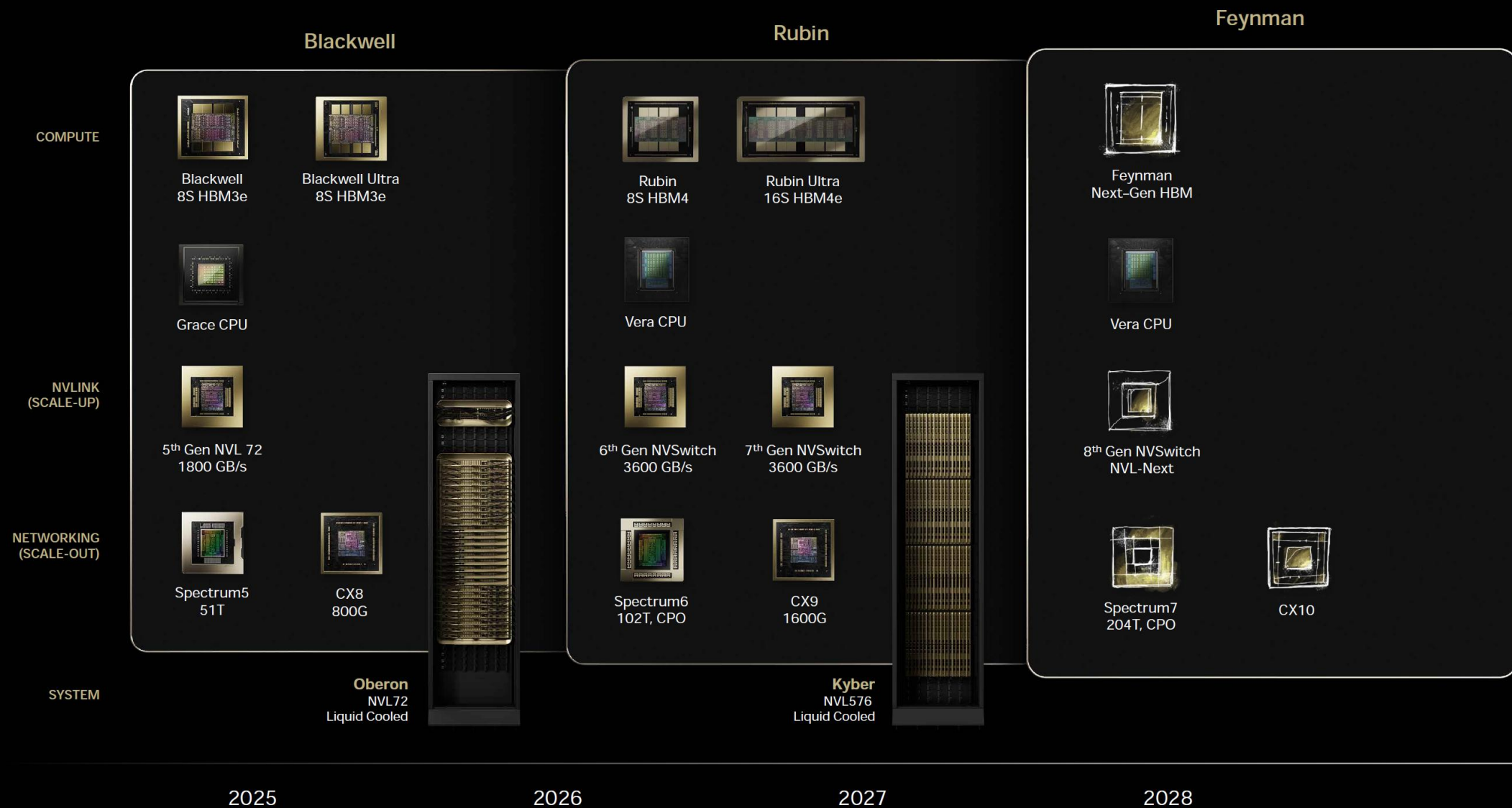


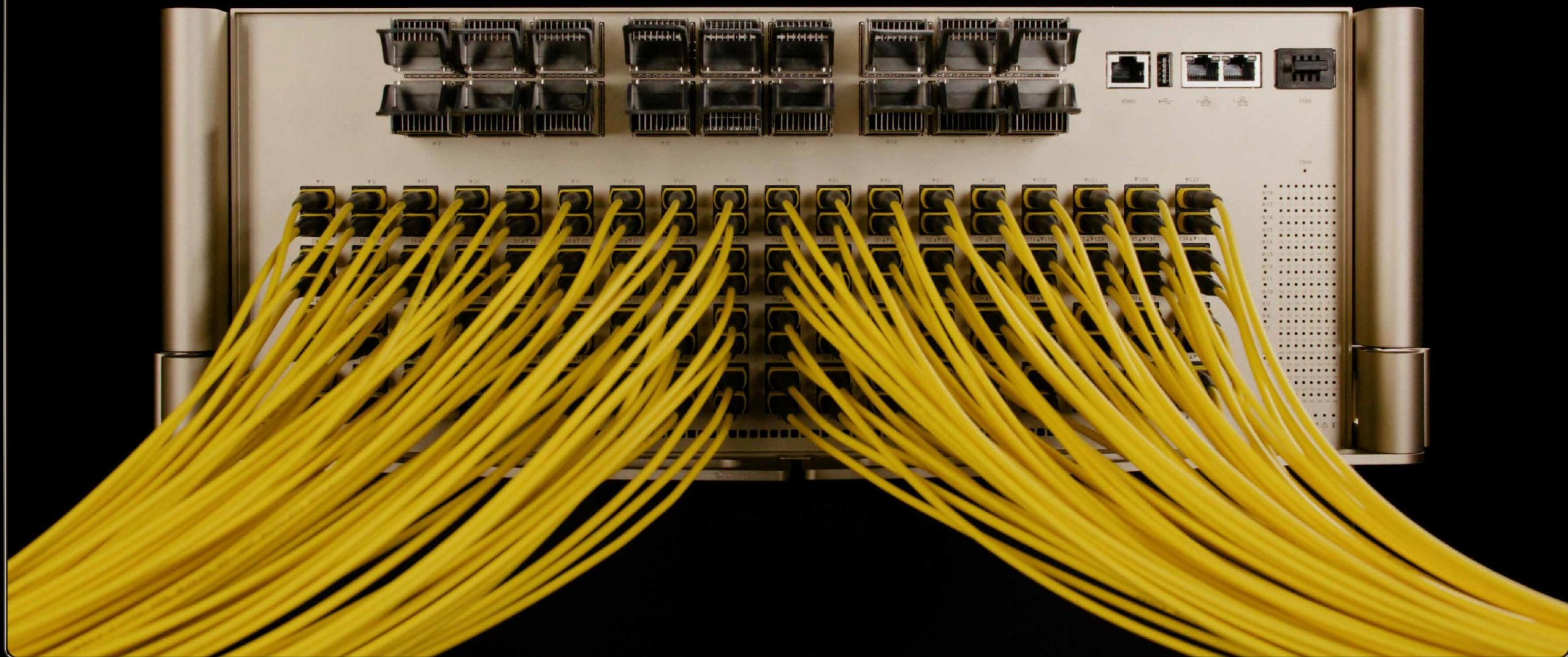
Source: Datadog

[출처] 10 Insights on real-world container use, Nov 2023 Datadog

NVIDIA Paves Road to Gigawatt AI Factories

One-Year Rhythm | Full-Stack | One Architecture | CUDA Everywhere







TEN은,
효율적인 솔루션을 제시합니다.



TEN은 AI를 잘 개발하고 운영할 수 있도록,
MLOps 소프트웨어와 **AI 전용 인프라**를 서비스하고 있습니다.



AI 학습과 운영에서 인프라 활용 방식이 기존 IT 서비스와 다르다는 것이 AI 난제의 시작



모델 학습을 위한 GPU 인프라 문제 24 X 7, GPU 자원 가동률 극대화

- AI 학습의 계산량을 위해 다수의 GPU와 노드를 '묶어서' 사용 해야 함
- GPU와 노드를 '묶어' 쓰기 위해서는 네트워크와 스토리지 역할이 중요

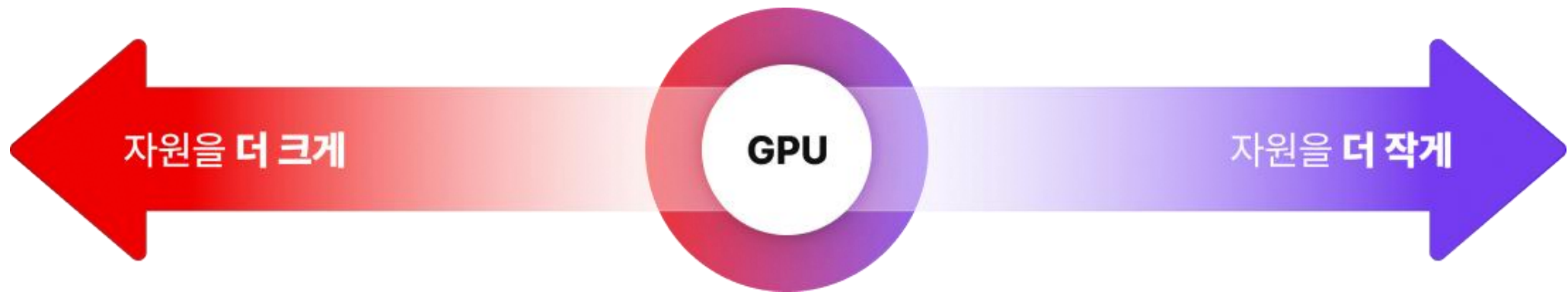


모델 운영을 위한 GPU 인프라 문제 최소한의 GPU 사용으로 비용절감과 운영 품질 확보

- AI 추론의 계산량은 적은 GPU 자원으로도 처리 가능
- AI 비즈니스 전개 시 인프라 비용 최소화가 사업의 성패를 가름
AI 서비스 원가로서 장기적으로 영향을 미침

생산성

저비용



차세대 AI 워크로드 오케스트레이션 플랫폼 AI Pub의 제품 라인업

K8s 관리를 위한



AI Pub K10s

AI 개발을 위한



AI Pub Dev

AI 운영을 위한



AI Pub Ops

AI 인프라 구축을 위한
컨설팅 서비스



RA:X

for

Kubernetes에 능숙한 개발자
DevOps 및 MLOps 개발자

Kubernetes 오브젝트 관리

for

Kubernetes에 익숙하지 않은
AI 모델러 혹은 운영 관리자

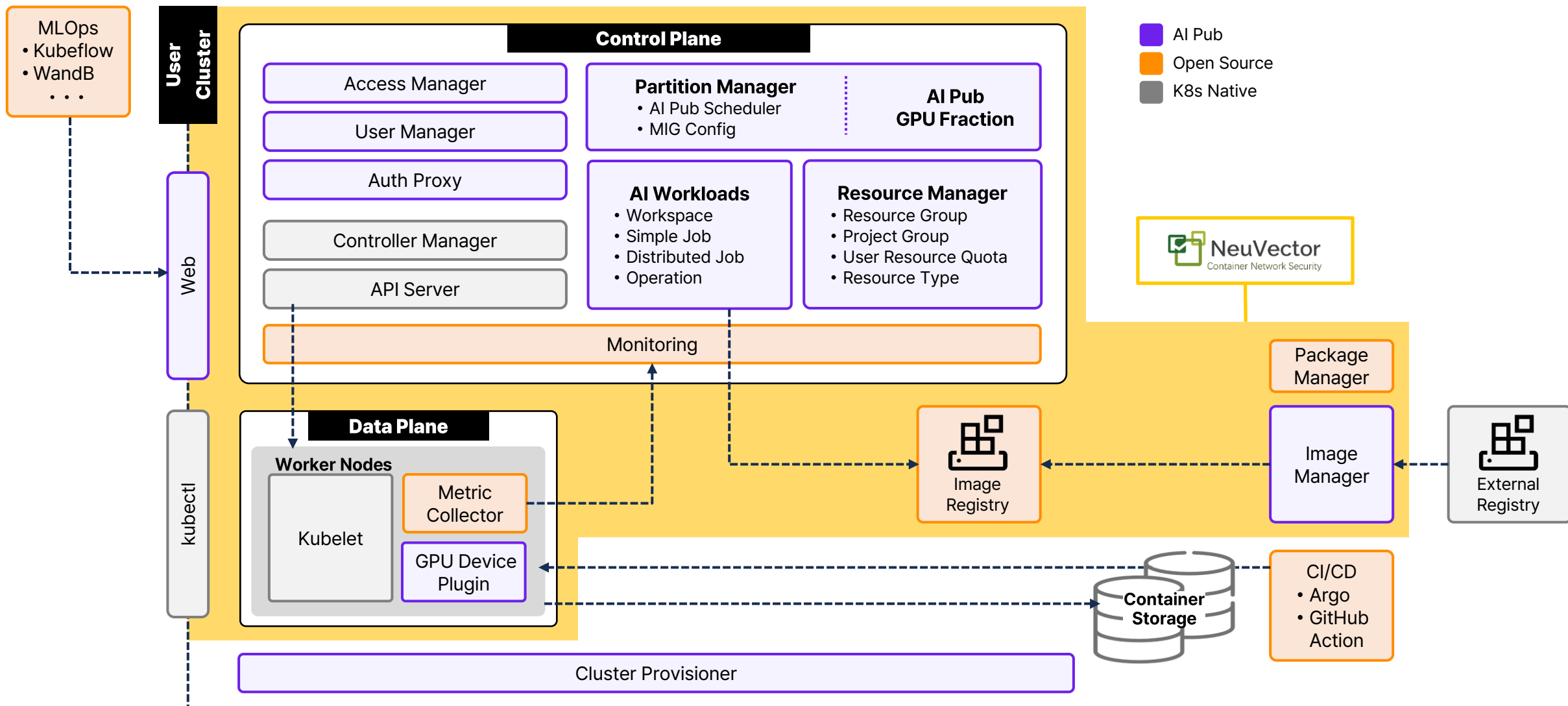
인공지능 학습 오브젝트 관리

for

AI인프라 신규 도입
또는 운영중인 기업

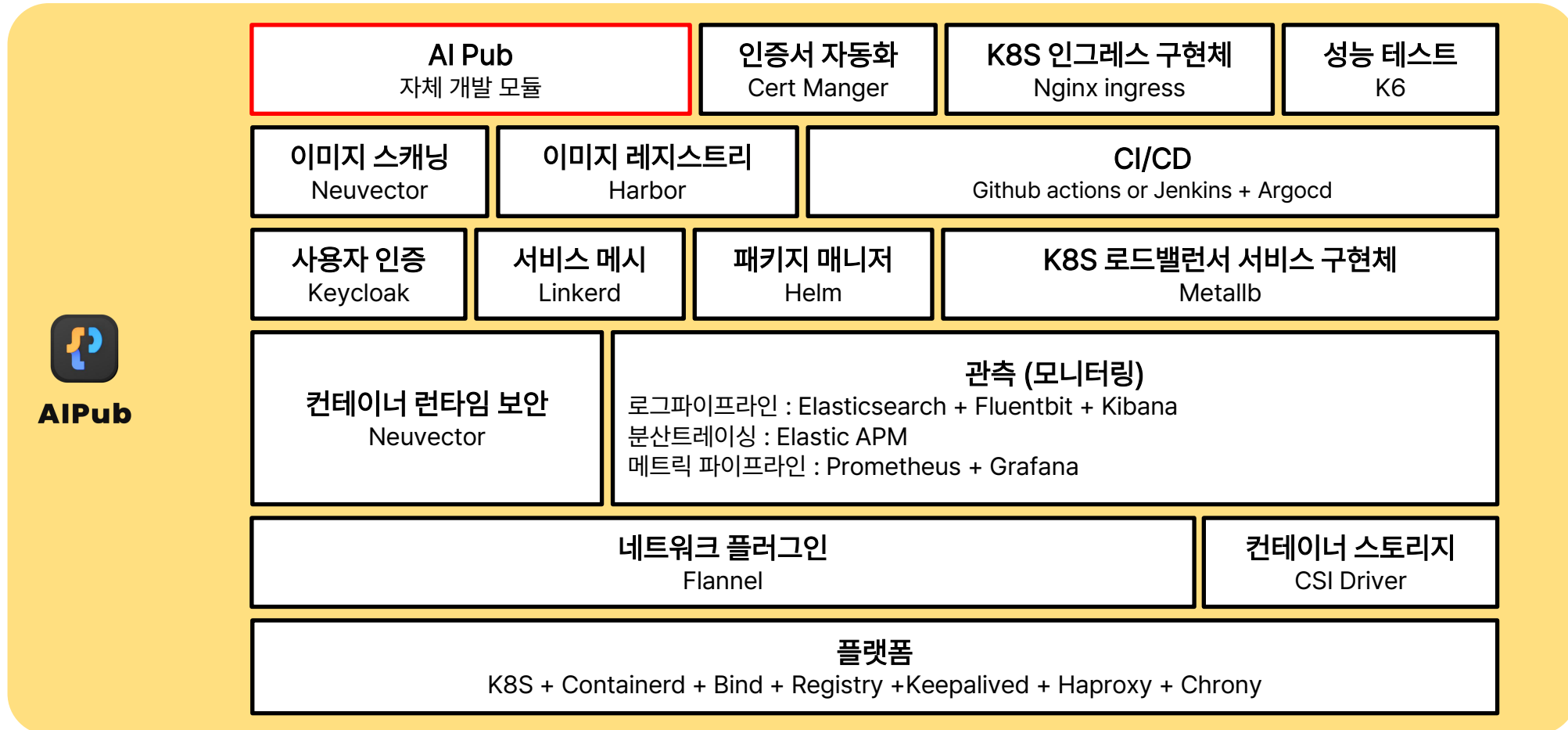
비용 효율적인 맞춤형 인프라 구성

Kubernetes 기반 AI Pub 플랫폼 구조



3. Kubernetes-Native Architecture

Kubernetes 에코 시스템의 검증된 도구들을 연동하여 서비스를 구현



GPU 활용률 극대화로 자원 낭비 최소화 및 비용 절감

복잡한 환경 설정 없이 인프라 구축 시간을 획기적으로 단축하며, 최적의 AI 성능 및 비용 효율성 동시 확보



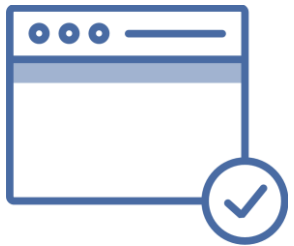
Ready-to-Accelerate

복잡한 환경 구성 없이 바로 사용 가능한 플랫폼으로 인프라 세팅과 AI 배포 기간을 획기적으로 단축



Maximized Efficiency

학습 단계에서는 GPU 자원을 최대한 활용하고, 운영 단계에서는 자원 사용을 최소화하여 비용은 줄이고 성능은 극대화



Faster to Market

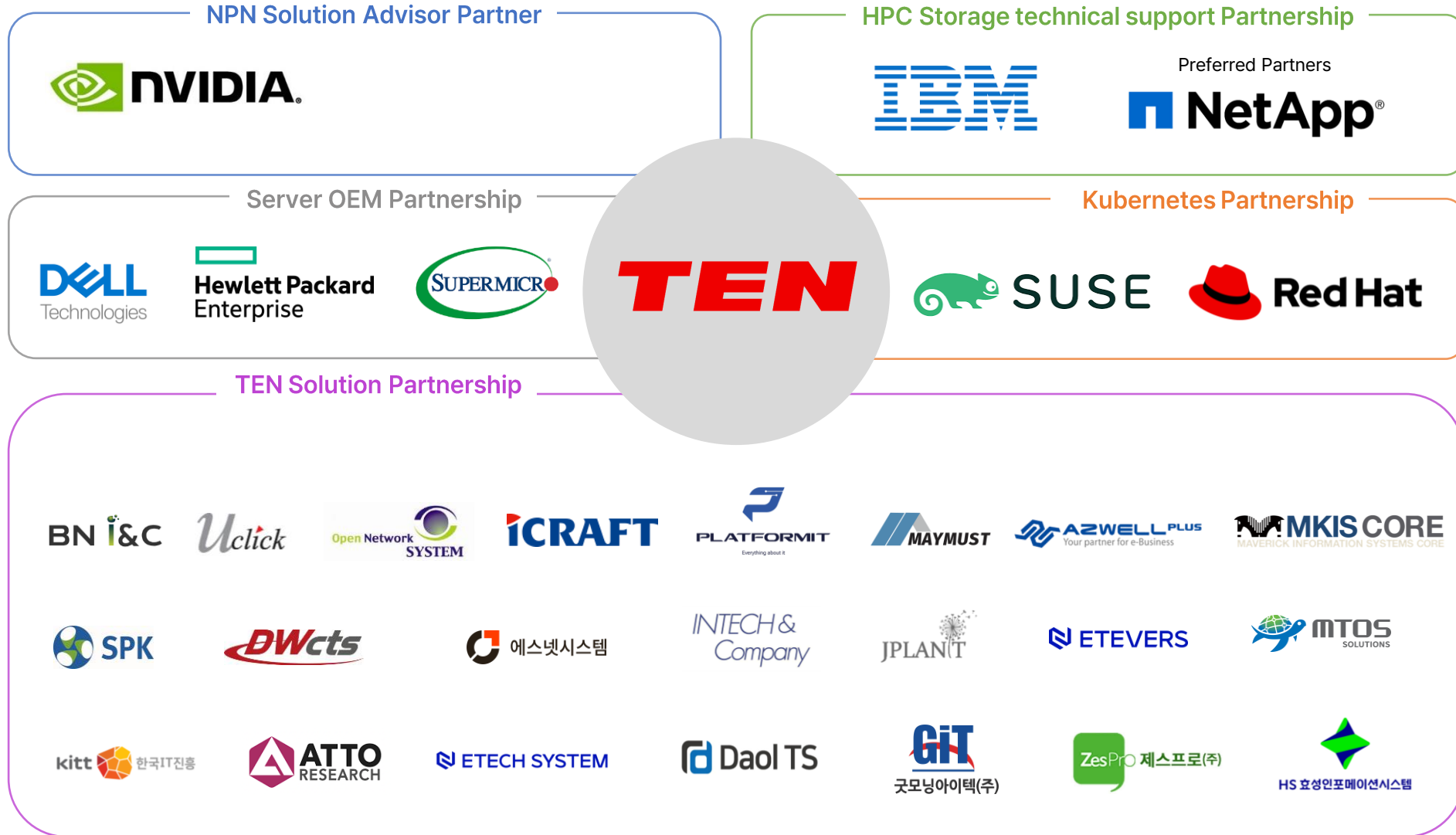
복수 워크로드, 사용자의 오케스트레이션
→ 빠른 반복과 짧은 배포 주기를 실현



Control & Visibility

직관적인 사용자 인터페이스와 실시간 리소스 모니터링 기능을 통해, 인프라 상태를 파악하고 효율적으로 관리 가능

AI 전문 밴더들과 파트너십을 통해
AI 전용 서비스를 올인원으로 제공하고 있습니다.



다양한 분야의 기업 및 학교에서 **TEN**의 서비스를 이용하고 있습니다.

[Client] 정부 및 기관



[Client] 기업



[Client] 교육



MLOps 관련 업체 중 국내 최초 Solution Advisor NVIDIA의 공식 파트너로 기술력과 서비스 역량을 인증 받았으며, NVIDIA의 다양한 솔루션을 통합하여 서비스를 제공할 수 있습니다.



Solution Advisor

Partners whose primary business model is to provide consultation services and expert advice to customers looking to implement NVIDIA products, NVIDIA-based solutions, and technologies.



Resource Management

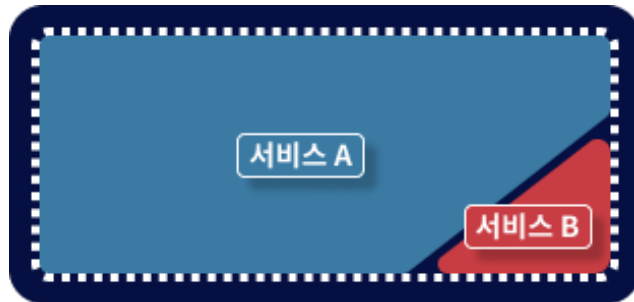
GPU와 기타 컴퓨팅 자원을 얼마나 효과적으로 분배하고
제어할 수 있는지를 확인합니다.

1. GPU 분할 가상화
2. 자원 관리 기능

1. GPU 분할 가상화

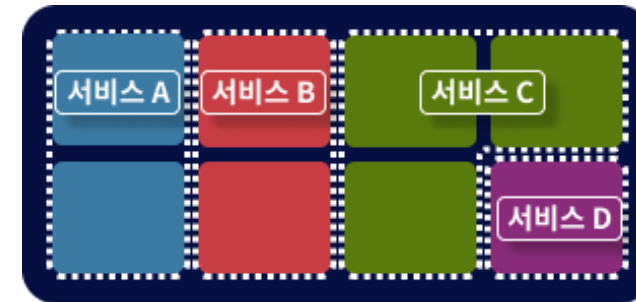
- AI Pub은 Kubernetes-native 방식으로 GPU 분할 기능을 제공 (NVIDIA Volta 이상 GPU를 지원)
- Device Plug-in 및 K8s Extended Resource와 통합되어 자원 할당의 유연성과 가시성 제고

쿠버네티스 기본 GPU 할당



- GPU는 1개 단위로만 요청 가능
- 경량 워크로드도 전체 GPU 점유 → 자원 낭비
- 여러 서비스 간 공유 불가능
- 자원 충돌 제어 어려움

AI Pub의 GPU 분할 할당 방식



1개 GPU에 여러 개의 서비스 운영이 가능
서비스간 리소스 사용 침해를 막아 안정성 확보

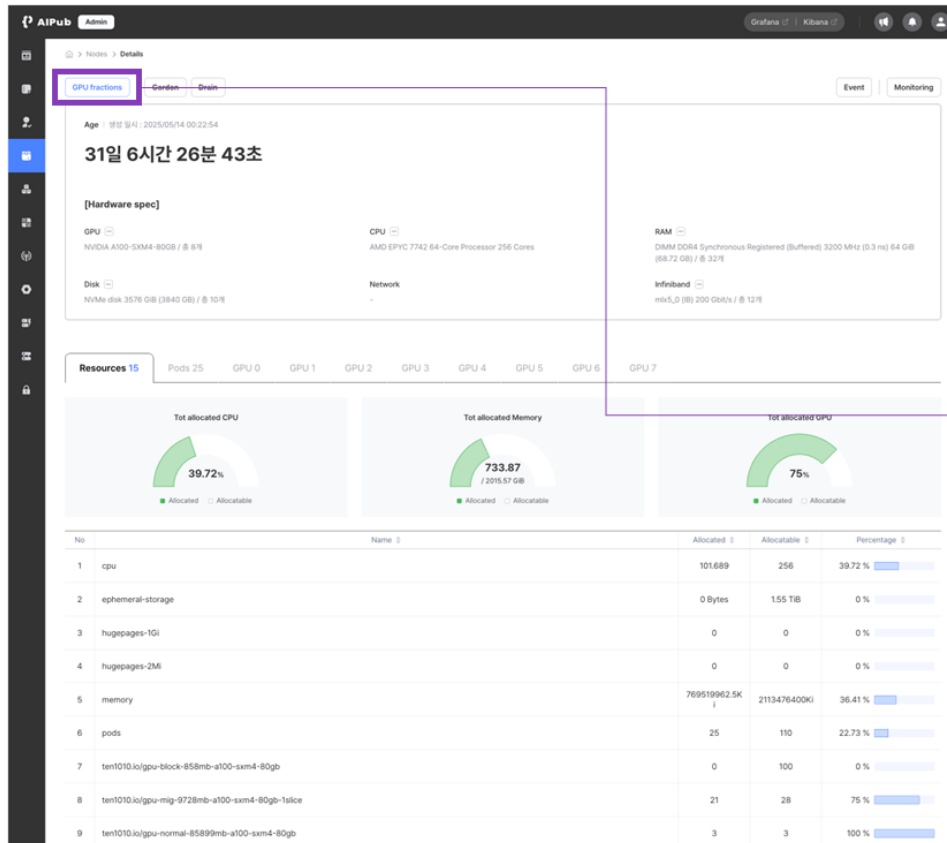
- GPU 1개의 코어와 메모리를 1% 단위로 분할, 100개 Block으로 분할
- 하나의 GPU를 여러 서비스가 동시에 사용
- Kubernetes 기반 리소스 스케줄링 가능
- 서비스 간 자원 침해 없이 안정적 운영

1. GPU 분할 가상화

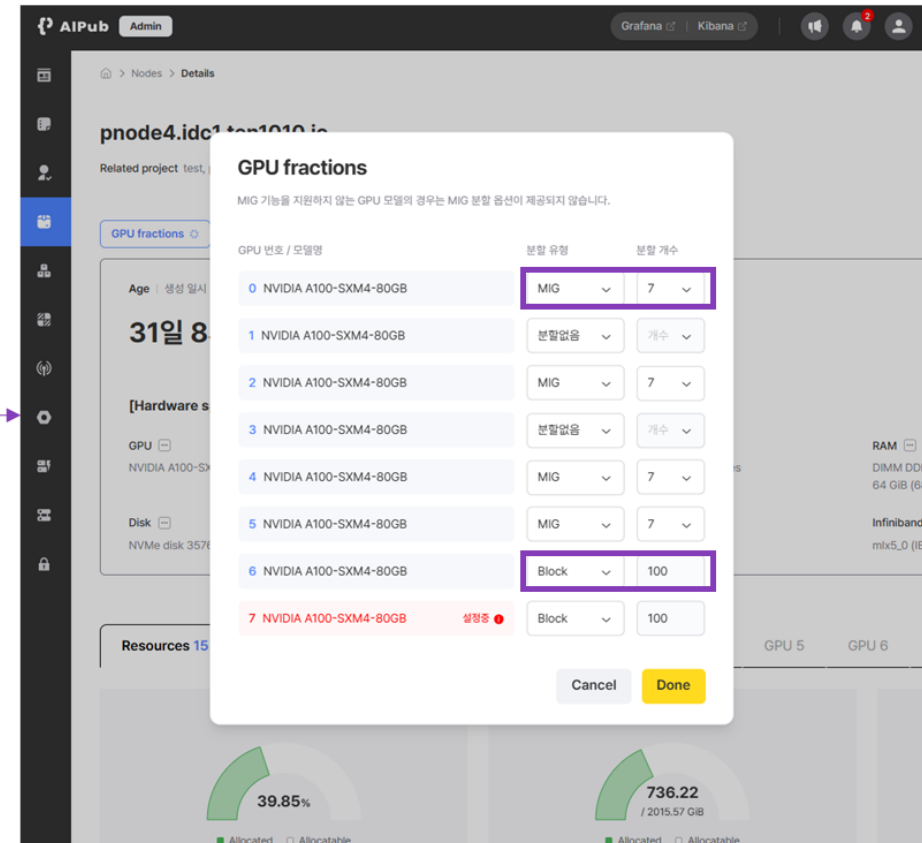
• 관리자 지정 GPU 분할 용이성

- GPU 별로 UI를 통해 분할 설정 가능: 분할 타입 (분할 없음 / MIG 분할 / 블록 분할) 및 분할 개수 설정 가능
- MIG 분할 설정은 MIG 분할 가능한 GPU인 경우에 한하여 가능

[노드 상세 화면 > Resources 탭]

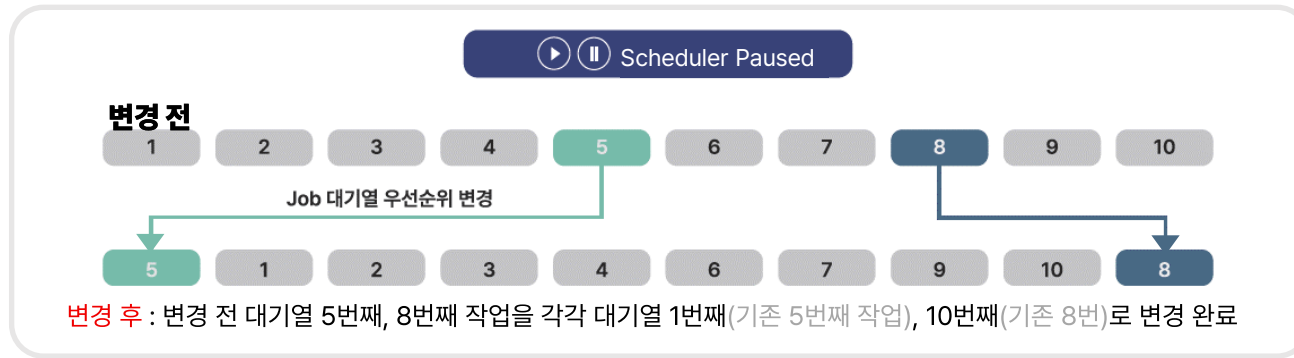


[GPU fractions 화면]



2. 자원관리 기능 - AI 워크로드에 최적화된 AI Pub 스케줄러

[워크로드 우선순위 조정 기능 예시]



AIPub Admin

Dashboard Projects Users Nodes ImageHub **Prioritizations** ResourceTypes Usage Workload

Queue Prioritization

작업 상단의 스케줄러 멈춤 버튼을 클릭하여 스케줄러 멈춘 상태가 되어야 대기열 우선순위 편집이 가능합니다.

기본 리소스 그룹

우선순위 목록

	운영자	이름	신청일시	리소스 타입 / 개수	순서 변경	우선순위
☒	김영주 idididididididid	workspace-1	2024/00/00 00:00:00	cpu-resource, CPU:2.99 코어, MEM:2.98 GiB	↕	내리기
☒	오승희 idididididididid	workspace-2	2024/00/00 00:00:00	nvidia-s100-pcie-40gb-m1g-2g.10gb / 99개 Multi node	↕	내리기
☒	신세훈 idididididididid	workspace-3	2024/00/00 00:00:00	cpu-resource, CPU:2.99 코어, MEM:2.98 GiB	↕	내리기

대기 목록

	운영자	이름	신청일시	리소스 타입 / 개수	우선순위
☒	정재환 idididididididid	job-name	2024/00/00 00:00:00	cpu-resource, CPU:2.99 코어, MEM:2.98 GiB	올리기
☒	김지훈 idididididididid	workspace-name 2	2024/00/00 00:00:00	nvidia-s100-pcie-40gb-m1g-2g.10gb / 99개 Multi node	올리기
☒	박순혜 idididididididid	workspace-team	2024/00/00 00:00:00	cpu-resource, CPU:2.99 코어, MEM:2.98 GiB	올리기

2. 자원관리 기능 - 정책 기반 리소스 관리

사용 정책에 따른 자원 할당 제어 및 자동 회수

- 사용자 별, 프로젝트 별 자원 할당 제어

ResourceSet(GPU/CPU/RAM)에 대해 사용자 별, 프로젝트 별 사용한도 설정과 저장공간(Volume) 사용한도 설정 기능 제공

- 워크스페이스 회수 기능

- 관리자가 저 사용, 미사용 워크스페이스에 대한 강제 회수 조치 가능
- 워크스페이스 생성 후 관리자가 지정한 시간 지나면 자동 회수 가능

- 배치 작업의 실행과 반복 회수 설정 가능

배치 작업을 예약 실행하고, 작업 종료 후 자원을 자동 회수함
반복적인 연산 업무를 코드 없이 UI 로 손쉽게 자동화할 수 있음

Set GPU quota for 홍길동

설정된 GPU 한도 내에서만 리소스를 사용할 수 있습니다. 기존에 설정된 Quota가 있는 User의 경우 최신 설정값으로 변경됩니다.

! 설정된 리소스 한도를 초과한 워크로드는 생성되지 않으며, 운영 중인 워크로드는 삭제될 수 있으니 설정에 유의해 주세요.

GPU 사용 한도 설정

ten1010.io/gpu-A100-1	1 개	🗑️
ten1010.io/gpu-nvidia-a100-pcie-40gb-mig-2g...	10 개	🗑️
ten1010.io/block-100	10 개	+

Cancel Done

Set workspace reclaim

이미 운영 중인 Workspace는 reclaim 설정 시점부터, 신규 Workspace는 리소스 할당 시점부터 회수 시간이 계산됩니다. 기존에 설정된 자동 회수 설정이 있는 User의 경우 최신 설정값으로 변경됩니다.

홍길동 / IDIDID × 홍길동 / IDIDID × 홍길동 / IDIDID × 홍길동 / IDIDID × 홍길동 / IDIDID ×

자동 회수 시간 설정

MAX 99999 시간 후 회수

Cancel Done

반복 주기 설정

섹션별 설명 내용이 나올 영역으로 최대 2줄을 표시할 예정 섹션별 설명 내용이 나올 영역으로 최대 2줄을 표시할 예정 섹션별 설명 내용이 나올 영역으로 최대 2줄을 표시할 예정 섹션별 설명 내용이 나올 영역으로 최대 2줄을 표시할 예정

주기

매주

세부 주기

월요일, 수요일

세부 실행 시각

자정

매시 정각

직접 입력

10시

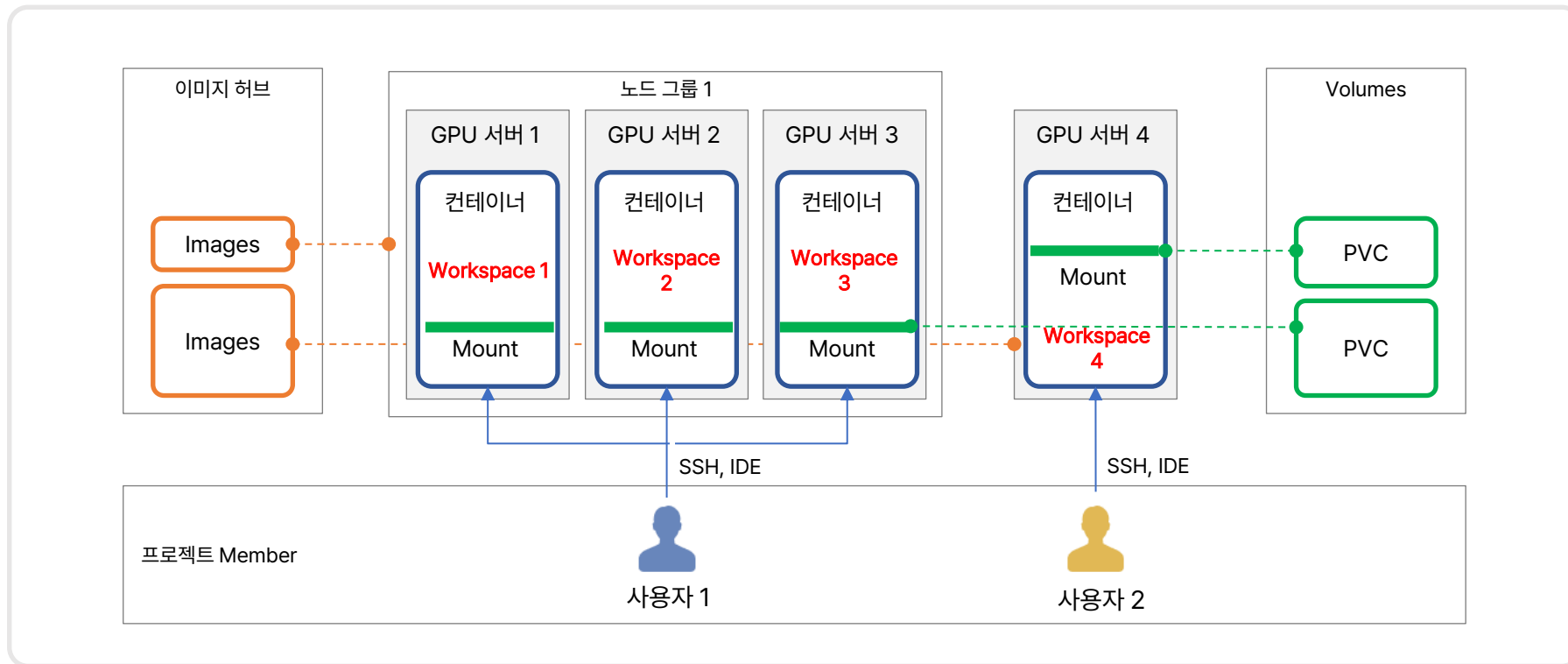
30분

설정 주기 Asia/Seoul 기준 매주 월요일, 수요일 10시 30분 마다 Job이 실행됩니다. (다음 실행 : 2025년 6월 1일 10시 30분)

2. 자원관리 기능- 프로젝트 기반 자원과 역할 관리

자체 개발한 '프로젝트' 컨트롤러를 통해 사용자와 자원을 유연하고 안전하게 매핑함
이를 통해 팀별, 목적별, 환경별로 논리적으로 격리된 리소스 그룹을 구성하고, 고도화된 자원 관리와 최적화를 지원

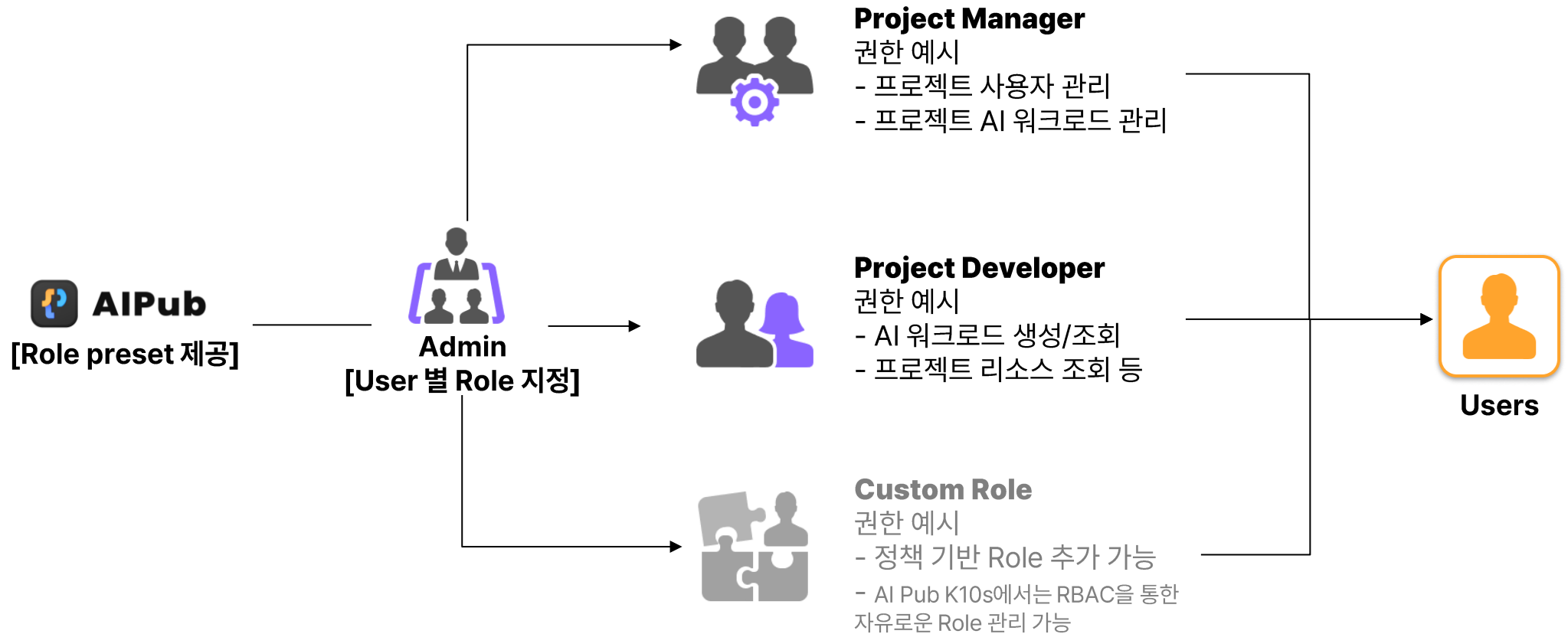
[프로젝트 별 리소스 접근 권한 관리]



2. 자원관리 기능- 프로젝트 기반 자원과 역할 관리

AI Pub은 역할 기반 사용자 권한 관리 기능을 제공함

쿠버네티스 RBAC 을 통해 사용자 정의 역할 및 권한 생성 가능하며, 주요 역할 프리셋 제공



User 관리

AI Pub은 사용자별 세밀한 리소스 및 권한 제어, GPU 및 워크스페이스 자원 관리, 직관적인 UI 기반의 통합 액션 기능을 통해 AI 클러스터 운영자가 안정적이고 효율적으로 사용자 관리를 수행할 수 있도록 지원합니다.

Users 76 Active 65 | Inactive 11 | Administrator 20

Buttons: Set admin, Approve, GPU quota, Workspace reclaim, Delete, User 전체, 검색어를 입력해 주세요.

No	ID	Name	Access	Creation time	Status	Reset PW
1	seunghye	seunghye seunghye test@ten1010.io	Project 2 ImageHub 0	2025/06/23 17:40:48	Active	Reset
2	error-test	test notification notification-test@ten1010.io	Project 0 ImageHub 0	2025/06/20 11:21:42	Active	Reset
3	yejikim	김 예지 yejikim1@ten1010.io	Project 0 ImageHub 0	2025/05/15 10:59:24	Active	Reset
4	seunghoalarm	알림 테스트 seunghoalarm@ten1010.io	Project 1 ImageHub 2	2025/05/13 21:46:50	Active	Reset
5	ritz85	최영 조 1111111@naver.com	Project 0 ImageHub 0	2025/05/09 15:58:19	Active	Reset
6	testtest111	test test21 wiehf@lkdjwio.com	Project 0 ImageHub 0	2025/04/26 19:24:17	Active	Reset
7	testuser9	생성 테스트 testuser9@ten1010.io	Project 0 ImageHub 0	2025/04/25 18:22:40	Active	Reset
8	happycat	asdf 한글1123asdf yunju.jung@ten1010.ionklsandflkandsnflk	Project 0 ImageHub 0	2025/04/25 18:21:36	Inactive	Reset
9	testuser8	초기화 정운주 testuser8@ten1010.io	Project 0 ImageHub 0	2025/04/25 17:45:54	Active	Reset
10	testuser7	잠금 정운주 testuser7@ten1010.io	Project 0 ImageHub 0	2025/04/25 17:29:07	Active	Reset

50 Items per page < 1 / 2 >

poadmin2

Manage account

GPU quota -

Name: 정운주, Email: poadmin2@ten1010.io, Phone: -, Department: -, Creation time: 2025/04/10 02:55:21, Comments: -

Project access 3

Name	Linked imageHub access	Roles
proj-yunju-2	yunju-imagehub-private / 권한 없음	
proj2	common / 권한 있음	Project developer
proj1	alpub-report / 권한 없음	Project manager

ImageHub access 2

Name	Roles
common	ImageHub developer

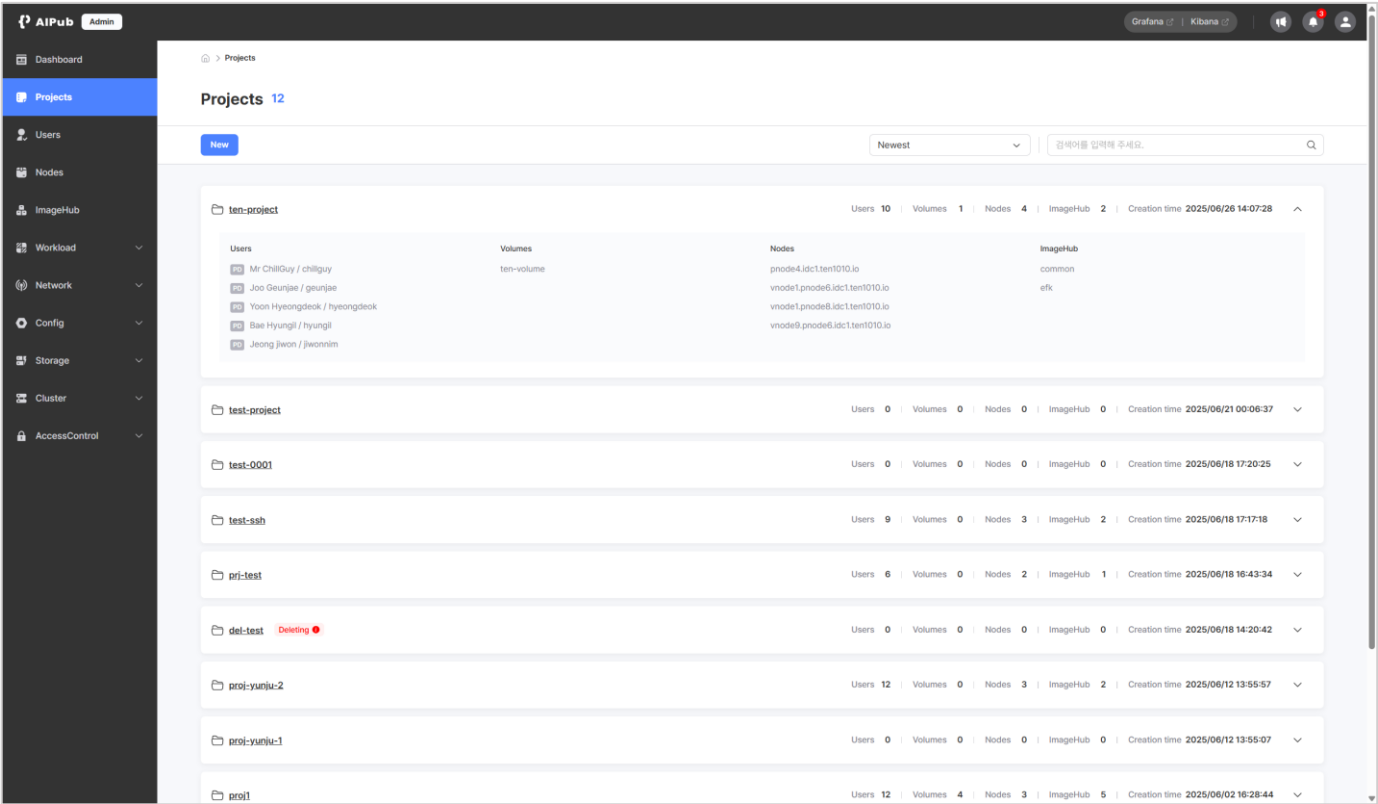
Linked imageHub access

- 1 common 권한 있음
- 2 dfsaf 권한 없음
- 3 seunghye-imagehub 권한 있음
- 4 yunju-imagehub-private.ac 권한 없음

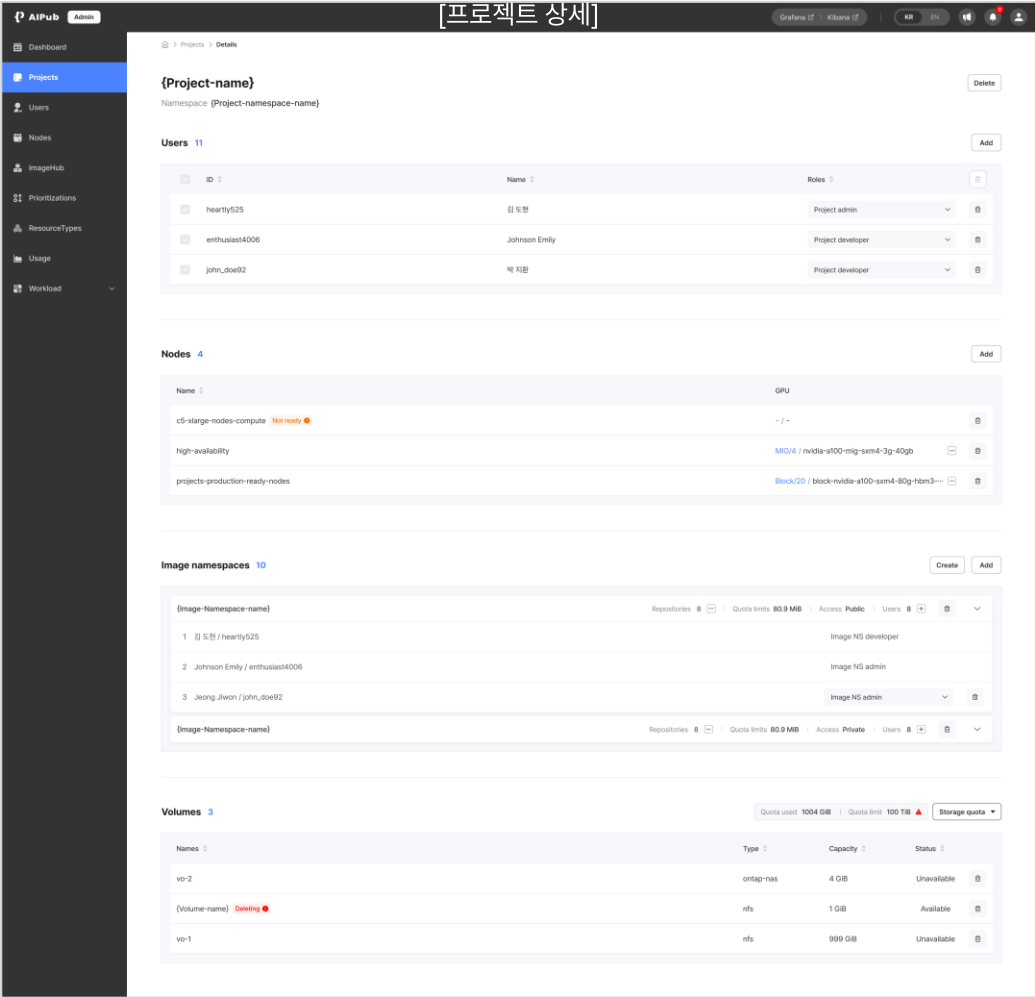
Project 관리

AI Pub의 프로젝트 관리 기능은 사용자, 볼륨, 노드, 이미지 저장소를 프로젝트 단위로 통합 관리하고 실시간 상태를 시각적으로 제공함으로써, 쿠버네티스 클러스터 운영자가 자원을 최적화하고 프로젝트별 협업을 안전하게 제어할 수 있도록 지원합니다.

[프로젝트 목록]



[프로젝트 상세]



Control and Visibility

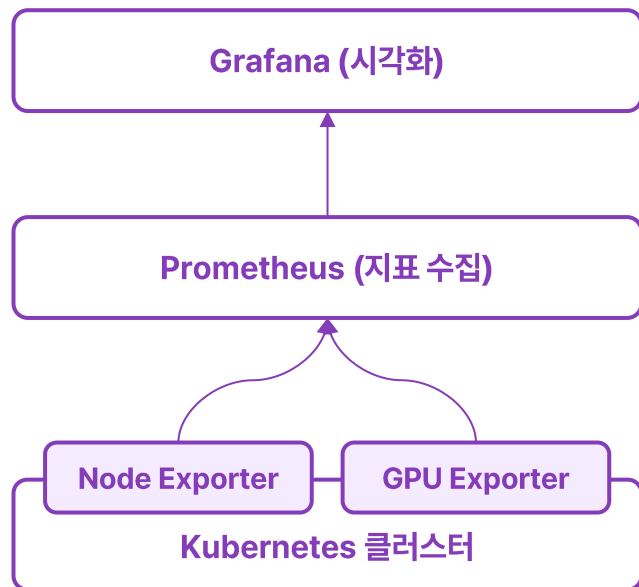
AI Pub의 서비스 별 UI와
모니터링 및 로그 기능을 소개합니다.

1. Monitoring – 메트릭 대시보드
2. Monitoring – 로그/이벤트 알림
3. 보고서 기능

1. AI Pub Monitoring: 메트릭 대시보드 개요

AI Pub은 리소스 사용량, 워크로드 상태, 성능 지표를 시각화한 대시보드를 통해 인프라와 서비스 상태를 한눈에 파악할 수 있음

[AI Pub의 메트릭 대시보드 스택 구성도]



1. 스택 개요

- AI Pub의 메트릭 대시보드 서비스는 Kube-Prometheus-Stack + Grafana 기반으로 구성되어 있음
- 인프라 및 워크로드 지표를 통합 수집하여 시각화 함

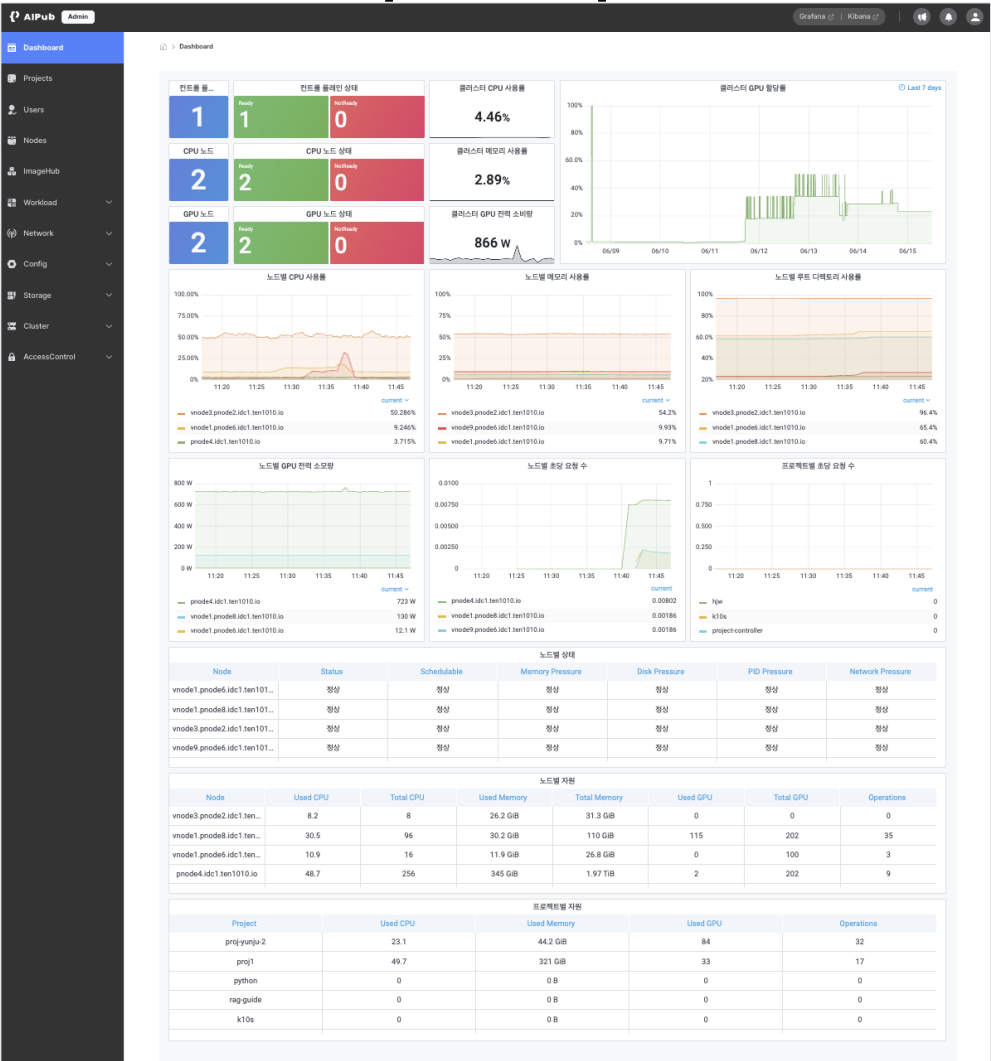
2. 커스텀 대시보드 구성 용이성

- Grafana 기반으로 고객 맞춤형 대시보드 구성이 자유로움
- Prometheus에서 수집되는 메트릭은 Kubernetes 오브젝트 단위로 구조화되어 있어 다양한 필터 및 정렬 가능
- 고객 요구에 맞는 지표 그룹핑 및 조건 기반 시각화(경고 임계값, 클러스터/네임스페이스 단위 등)가 가능
- 기존 템플릿 복제 및 변수 활용을 통해 비개발자도 손쉽게 대시보드를 편집가능

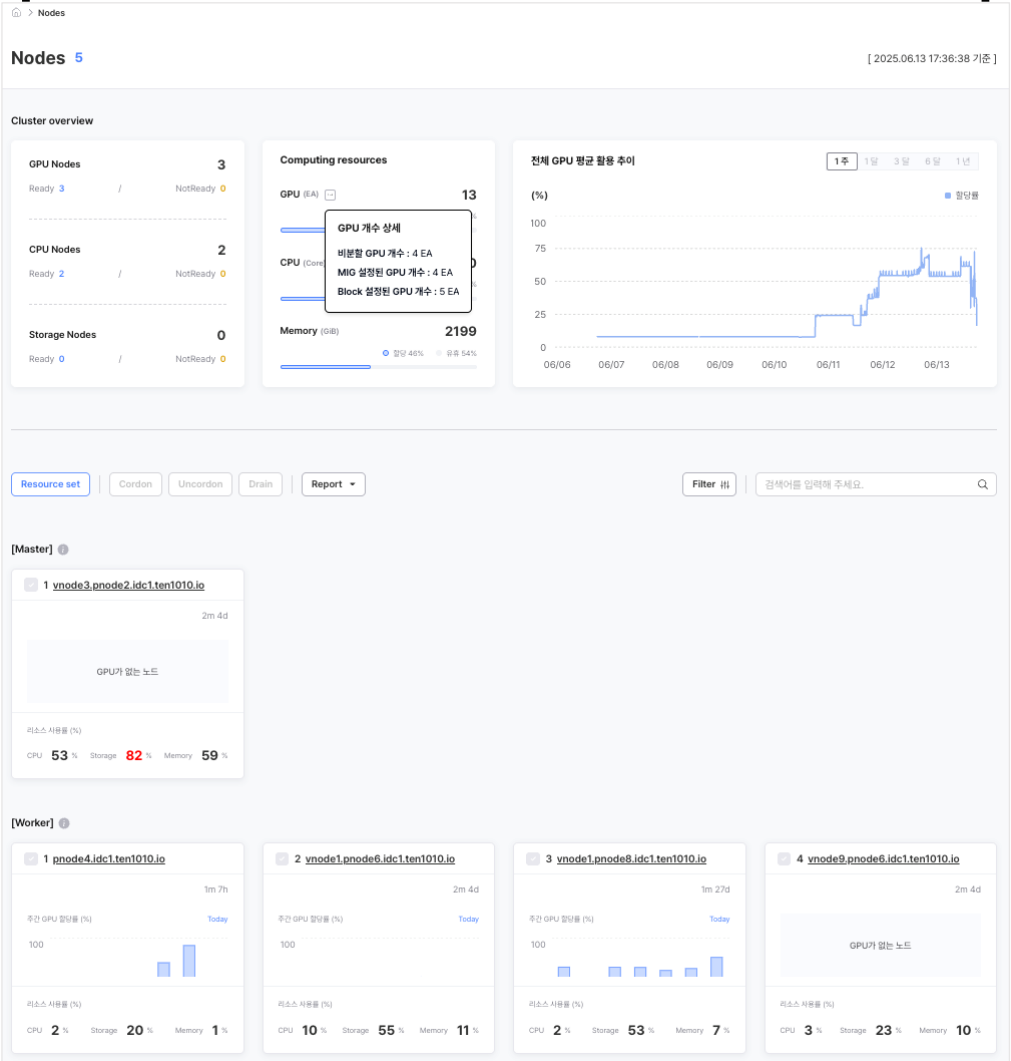
1. AI Pub Monitoring: 메트릭 대시보드

클러스터와 노드 모니터링: 클러스터 컴포넌트 별 통합/개별 지표 및 GPU의 통합/분할 설정 별 상세 지표 제공

[메인 대시보드]



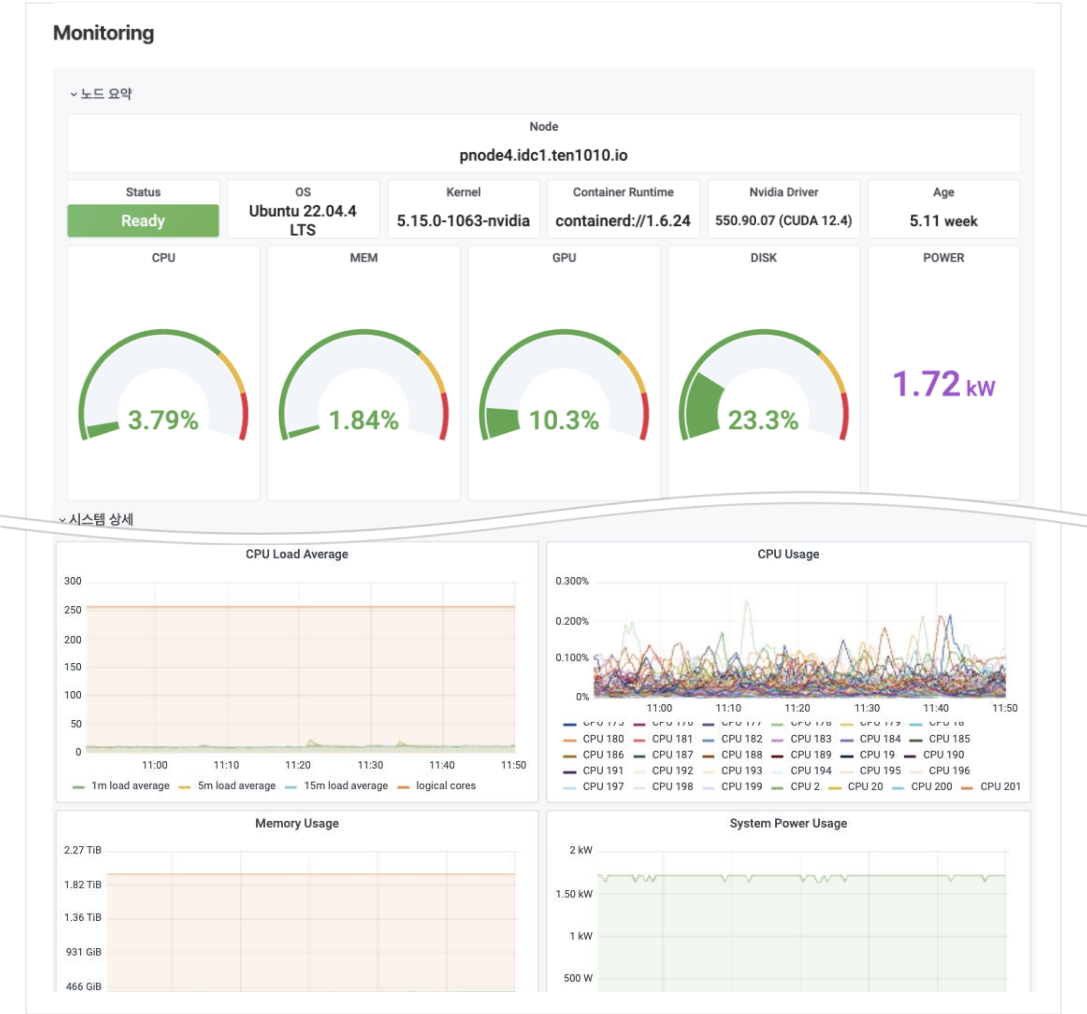
[노드 목록 화면-클러스터 요약 및 노드 별 리소스 사용 상태 시각화]



1. AI Pub Monitoring: 메트릭 대시보드

노드 별 모니터링: GPU 관련 다양한 메트릭은 노드별 모니터링 페이지에서 별도 대시보드로 구성 및 제공

[노드 별 모니터링 1/4 : 요약]



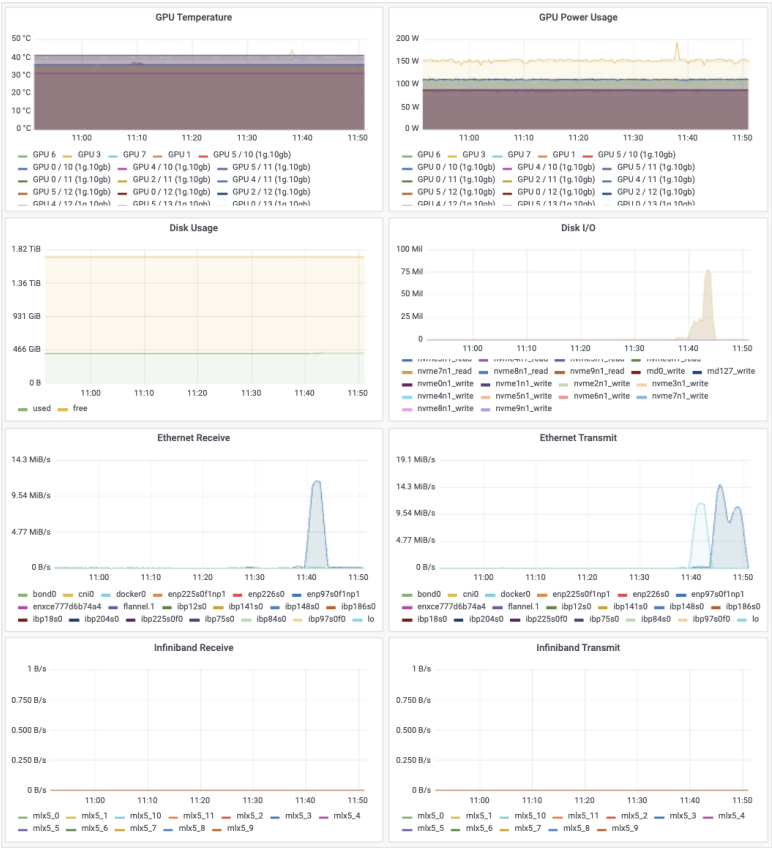
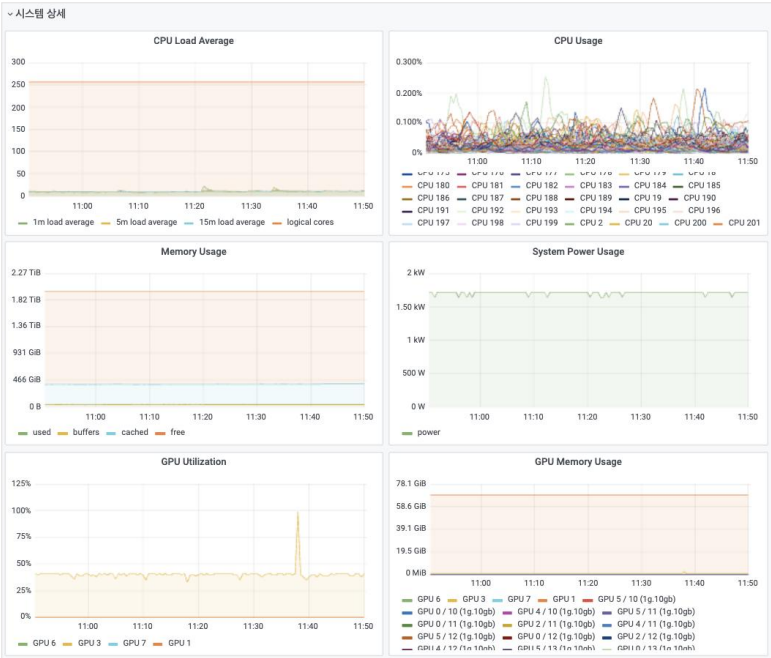
[노드 별 모니터링 1/4 : 요약 – 메트릭 상세]

패널	설명
Node	노드 이름
Status	노드 준비 상태 (Ready/Not Ready)
OS	운영체제 이름 및 버전
Kernel	커널 버전
Container Runtime	컨테이너 런타임 종류 및 버전
Age	생성 후 경과 시간
CPU	논리 코어별 사용률
MEM	전체 메모리 사용률
GPU	전체 GPU 사용률
DISK	디스크 사용률
POWER	전력 소비량

1. AI Pub Monitoring: 메트릭 대시보드

노드 별 모니터링: GPU 관련 다양한 메트릭은 노드별 모니터링 페이지에서 별도 대시보드로 구성 및 제공

[노드 별 모니터링 2/4 : 시스템 상세]



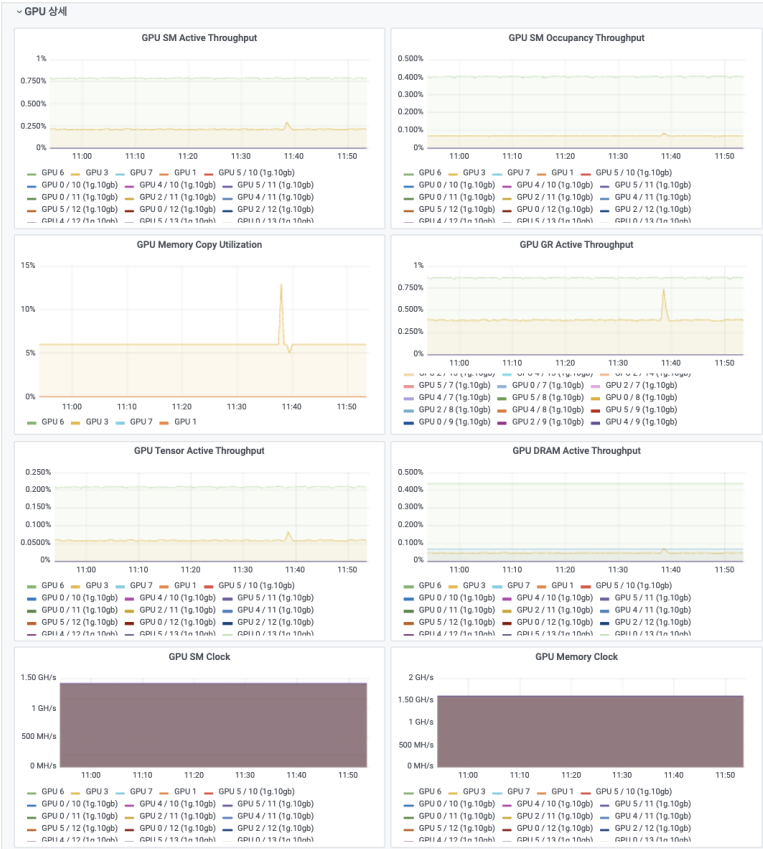
[노드 별 모니터링 2/4 : 시스템 상세 – 메트릭 상세]

패널	설명
CPU Load Average	1/5/15분 동안 평균사용한 논리 코어 수
CPU Usage	논리 CPU별 사용률
Memory Usage	메모리 사용량
System Power Usage	시스템 전체 전력 사용량
GPU Utilization	GPU 사용률
GPU Memory Usage	GPU 메모리 사용량
GPU Temperature	GPU 온도
GPU Power Usage	GPU 전력 소비량
Disk Usage	디스크 사용량
Disk I/O	디스크 읽기/쓰기 트래픽
Ethernet Receive	네트워크 수신 트래픽
Ethernet Transmit	네트워크 송신 트래픽
Infiniband Receive	인피니밴드 수신 트래픽
Infiniband Transmit	인피니밴드

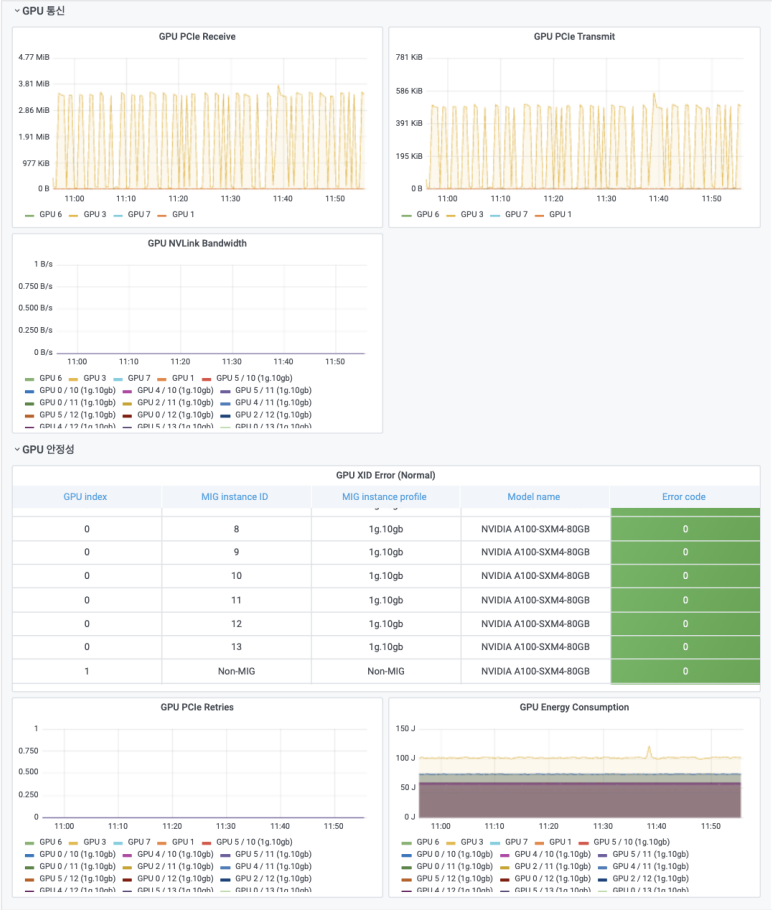
1. AI Pub Monitoring: 메트릭 대시보드

노드 별 모니터링: GPU 관련 다양한 메트릭은 노드별 모니터링 페이지에서 별도 대시보드로 구성 및 제공

[노드 별 모니터링 3/4 : GPU 상세]



[노드 별 모니터링 4/4 : GPU 통신 및 안정성]



[노드 별 모니터링 3/4 : GPU 상세 – 메트릭 상세]

패널	설명
GPU SM Active Throughput	SM 유닛의 실행률
GPU SM Occupancy Throughput	SM 유닛 점유율
GPU Memory Copy Utilization	메모리 복사 활용률
GPU GR Active Throughput	그래픽 유닛 실행률
GPU Tensor Active Throughput	텐서 코어 실행률
GPU DRAM Active Throughput	DRAM 활용률
GPU SM Clock	SM 유닛 클럭 속도
GPU Memory Clock	메모리 클럭 속도

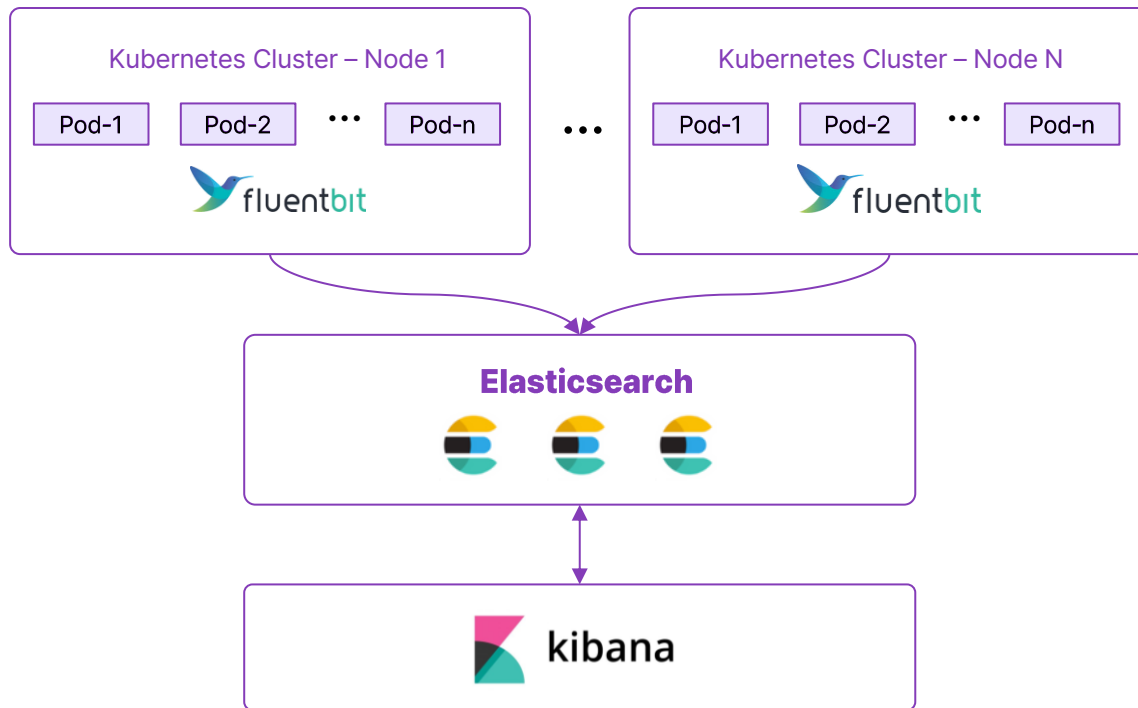
[노드 별 모니터링 4/4 : GPU 상세 – 메트릭 상세]

패널	설명
GPU PCIe Receive	GPU별 PCIe 수신 트래픽
GPU PCIe Transmit	GPU별 PCIe 송신 트래픽
GPU NVLink Bandwidth	GPU 간 NVLink 대역폭
GPU XID Error	GPU 오류 코드
GPU PCIe Retries	PCIe 재전송 발생률
GPU Energy Consumption	총 에너지 사용량 (joules)

2. AI Pub Monitoring: 로그/이벤트 알림 개요

에러, 이벤트, 사용자 행동 로그를 수집하여 실시간 알림과 함께 상세한 문제 원인 분석 기능을 제공

[AI Pub의 로그/이벤트 알림 스택 구성도]

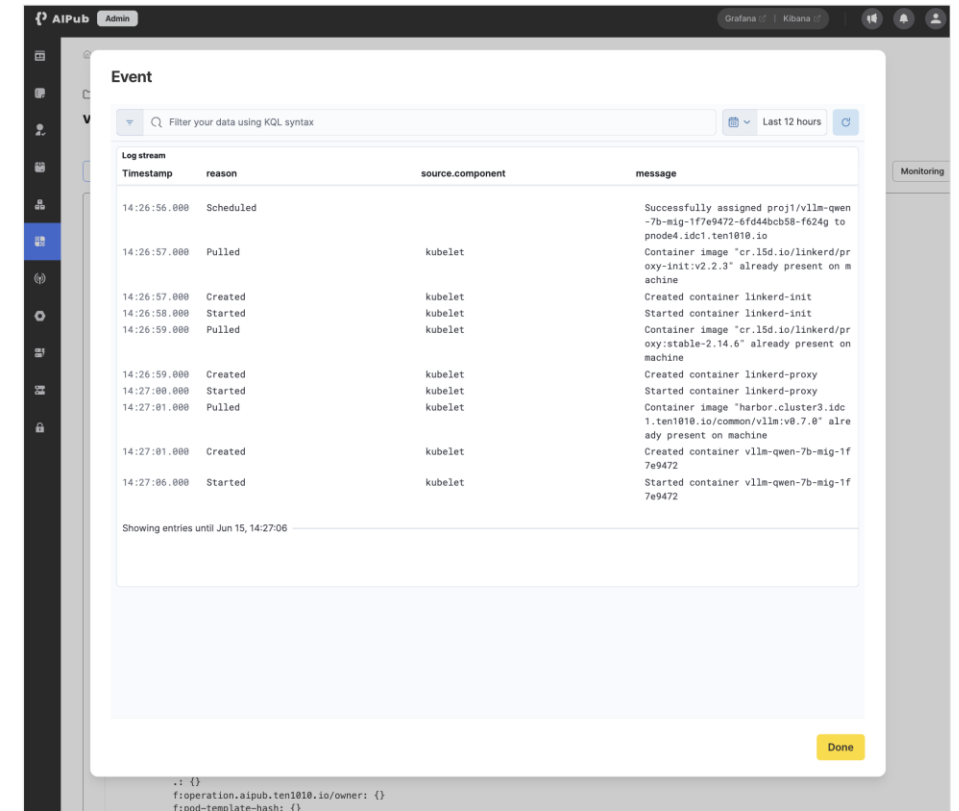
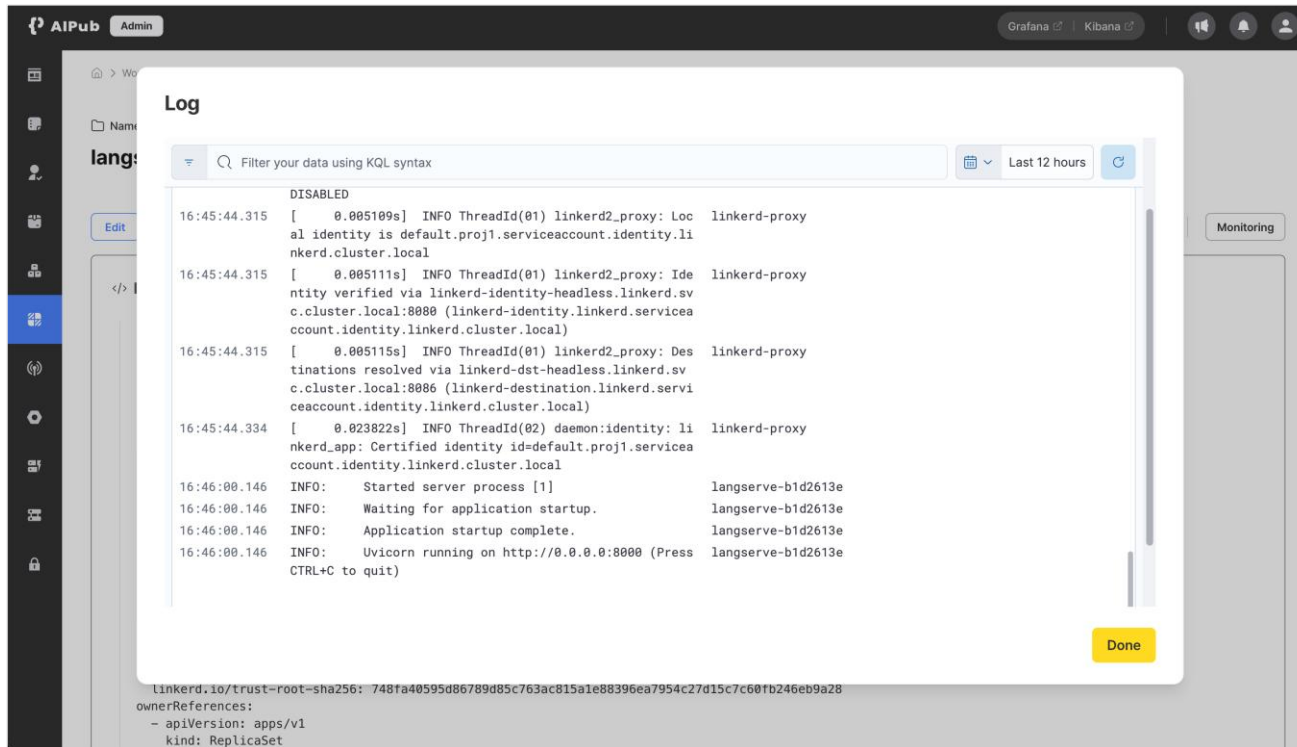


1. 스택 개요

- AI Pub의 로그/이벤트 알림 기능은 EFK(Elasticsearch, Fluent Bit, Kibana) 스택을 기반으로 구성되어 있음
- Kubernetes 기반 워크로드에서 발생하는 로그 및 이벤트 데이터를 수집/저장/시각화하는 구조로 구성
- 저장된 데이터는 검색 성능과 확장성을 보장하면서 Kibana를 통해 손쉽게 시각화 가능
- 수집된 데이터는 Elasticsearch를 통해 마스터 노드에 분산 저장하여 고가용성과 확장성을 동시에 확보

2. AI Pub Monitoring: 로그/이벤트 알림 개요

에러, 이벤트, 사용자 행동 로그를 수집하여 실시간 알림과 함께 상세한 문제 원인 분석 기능을 제공



3. AI Pub Monitoring: 보고서 기능

클러스터 현황에 대한 Daily/Weekly/Monthly PDF 리포트 제공 및 리소스 사용 내역 관리 기능

일일 노드 리포트

보고서 날짜: 2025년 06월 12일 (1)

서버 일일 성능 리포트

노드별 성능 현황

노드

vnodel.pnode6.idc1.te

pnode4.idc1.ten10i

vnodel.pnode6.idc1.te

vnodel.pnode8.idc1.te

vnodel.pnode8.idc1.te

리소스 사용률 비교

주간 메트릭 리포트

보고 기간: 2025-06-05 ~ 2025-06-12

주간 리소스 사용률 요약

CPU 사용률 일별 활용량

노드

pnodel.idc1.ten1010.io

vnodel.pnode6.idc1.ten1010.io

vnodel.pnode8.idc1.ten1010.io

vnodel.pnode8.idc1.ten1010.io

vnodel.pnode8.idc1.ten1010.io

vnodel.pnode8.idc1.ten1010.io

월간 노드 성능 리포트

보고 기간: 2025년 05월 13일 ~ 2025년 06월 12일 생성 시간: 2025-06-12 01:34:12

데이터 수집 방식 - 일별 평균: 하루 24시간 전체의 평균값 (5분 간격 샘플링) - 일별 최대: 하루 중 가장 높은 값 - GPU 온도: 노드별 모든 GPU의 평균(최대 온도 - GPU 에너지: 노드별 모든 GPU의 총 에너지 소비량 (mW → kWh) 변환) - 카디널리티 최적화: 데이터 매트릭 단순화로 성능 향상

데이터 수집 현황

메트릭	노드 수	데이터 포인트	카디널리티	상태
CPU 평균 사용률	5	154/150	정상	양호
CPU 최대 사용률	5	154/150	정상	양호
메모리 평균 사용률	5	154/150	정상	양호
메모리 최대 사용률	5	154/150	정상	양호
디스크 평균 사용률	5	154/150	정상	양호
디스크 최대 사용률	5	154/150	정상	양호
CPU 부하 온도	3	64/90	정상	소부분
CPU 최대 온도	3	64/90	정상	소부분
GPU 일일 에너지 소비	3	64/90	정상	소부분

데이터 수집 분석

- 데이터 누락 확인: Prometheus 보론 설정, 노드 다운타임, 네트워크 문제 등
- 카디널리티 관리: 불필요한 레이블 제거, 메트릭 집계 활용
- 장해 최적화: 수집 간격 조정, 메트릭 선택 수집
- 현재 상태: 수집 가능한 데이터로 분석 진행

성능 시각화 분석

클러스터 리소스 비교

모든 노드의 평균 리소스 사용률을 비교합니다. 노드간 부하 불균형을 확인하여 워크로드 재분배가 필요하지 판단할 수 있습니다.

Cluster Resource Utilization Comparison

Average CPU (%)

Average Memory (%)

Average GPU Temp (°C)

GPU 에너지 소비 트렌드

노드별 일일 GPU 에너지 소비량과 클러스터 전체의 에너지 소비 패턴을 분석합니다. 에너지 효율성 최적화에 활용될 수 있습니다.

Page 1 of 7

AI Pub Admin

스케줄러 방문 H

새로고침

자동 OFF

M

통합 알림 센터

GPU 및 서버 모니터링

이벤트 관리

Workloads 관리

대기열 우선 순위 관리

사용자 계정 관리

리소스 그룹 관리

사용 내역 관리

리소스 사용 내역

월별 리소스 사용 내역

2024/00/00 ~ 2024/00/00

적용하기

초기화

계정별 사용 내역

그룹별 사용 내역

2024년 12월

역설 다운로드

생성된 워크로드 개수

9,999 개

총 사용 시간

GPU 9,999 시간 59 분

CPU 9,999 시간 59 분

사용자 정보	이름	리소스 그룹	리소스 타입 / 개수	사용 기간	총 사용 시간
김나연	workspace1	resourceName1	nvidia-a100-8xm4-40gb / 2 개	2024/00/00 00:00:00 ~ 2024/00/00 00:00:00	9,999 시간 59 분
	workspace2	resourceName2	nvidia-a100-8xm4-40gb / 2 개	2024/00/00 00:00:00 ~ 2024/00/00 00:00:00	9,999 시간 59 분
	workspace3	기본 리소스 그룹	cpu-resource, CPU: 2.99 코어, MEM: 2.98 GiB	2024/00/00 00:00:00 ~ 2024/00/00 00:00:00 (가동중)	9,999 시간 59 분
	GPU 리소스 사용 소개		90 개		9,999 시간 59 분
CPU 리소스 사용 소개		-		9,999 시간 59 분	
박현우	workspace1	기본 리소스 그룹	nvidia-a100-8xm4-40gb / 2 개	2024/00/00 00:00:00 ~ 2024/00/00 00:00:00	9,999 시간 59 분
	workspace4	resourceName2	nvidia-8xm4-4000-nvidia-8xm4-a100 / 90 개	2024/00/00 00:00:00 ~ 2024/00/00 00:00:00 (가동중)	9,999 시간 59 분
	GPU 리소스 사용 소개		90 개		9,999 시간 59 분
	CPU 리소스 사용 소개		-		9,999 시간 59 분

TEN



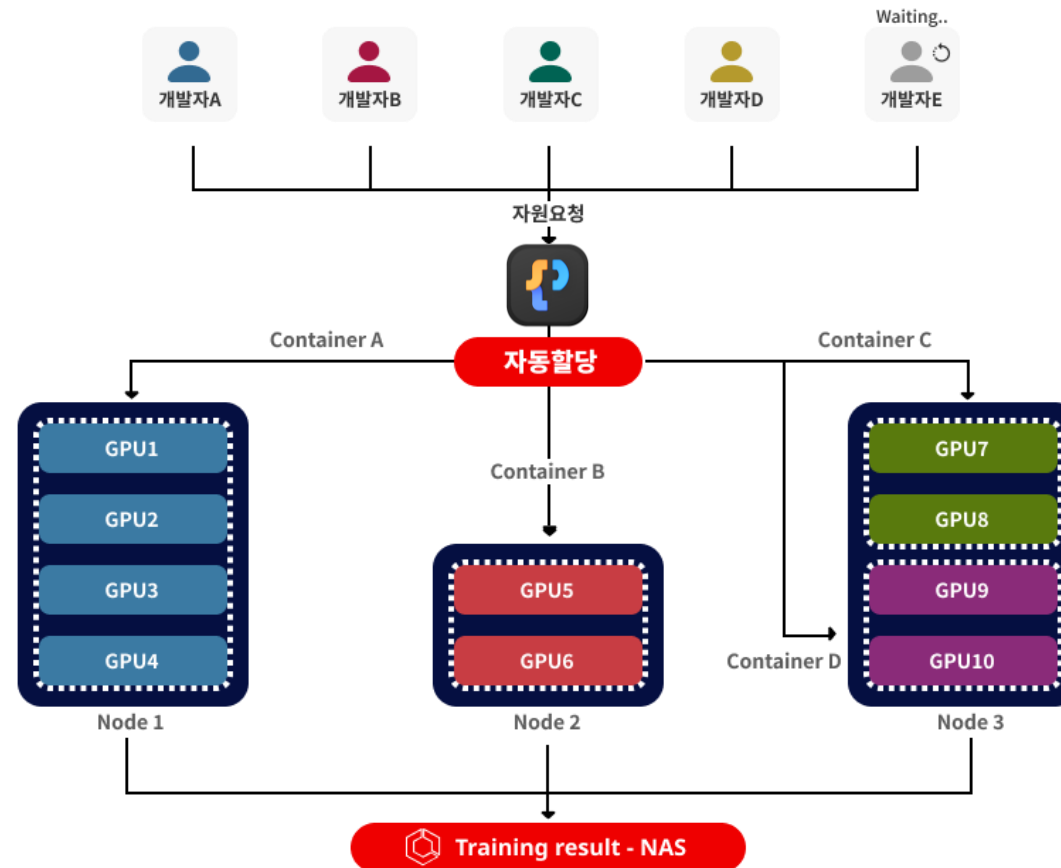
AI Pub Dev



COASTER를 코어로 하여 AI 개발을 지원하는 완전 관리형 서비스

팀 혹은 사용자 단위로 AI 인프라를 할당할 수 있습니다.
팀 개발자들의 자원 사용량 및 팀 별 사용량을 집계할 수 있습니다.

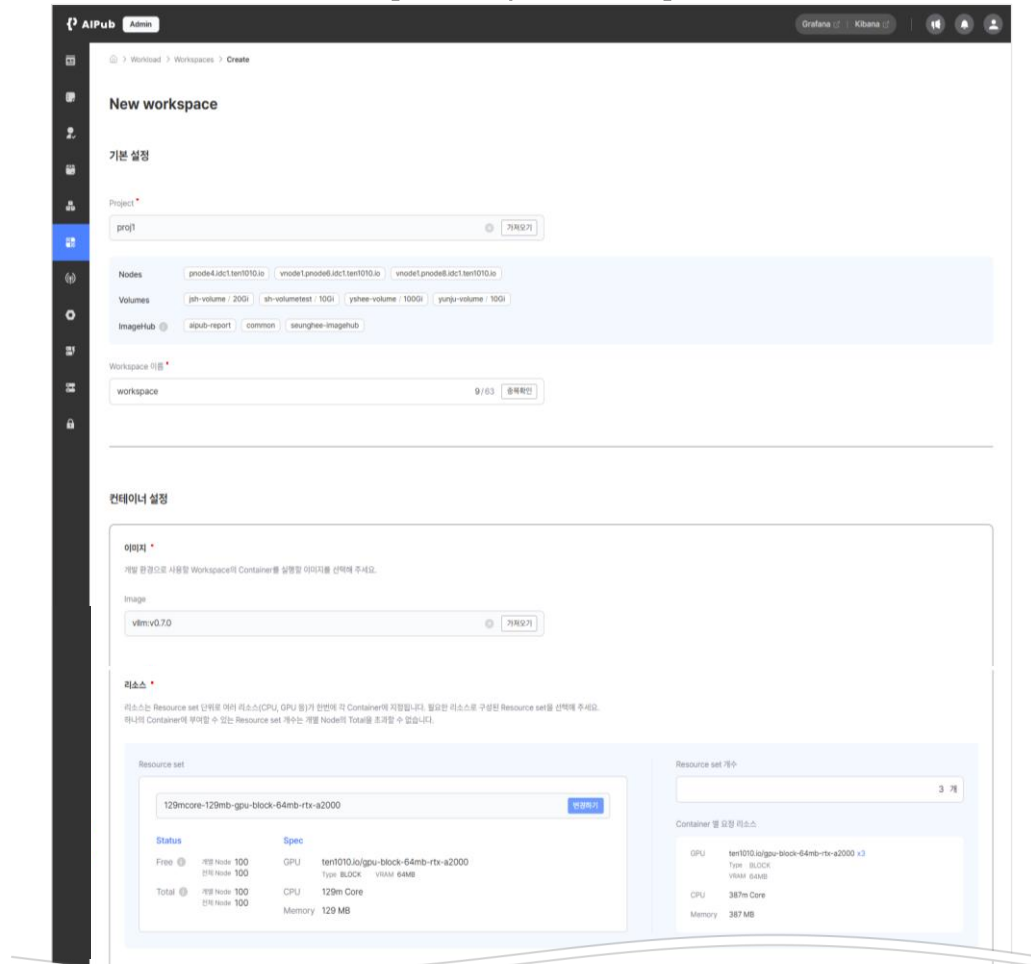
AI 개발용 인프라를 구축하여 중앙 관리할 때에도 자원 할당 및 관리가 가능하고, 각 단위 별 사용량을 측정할 수 있습니다.



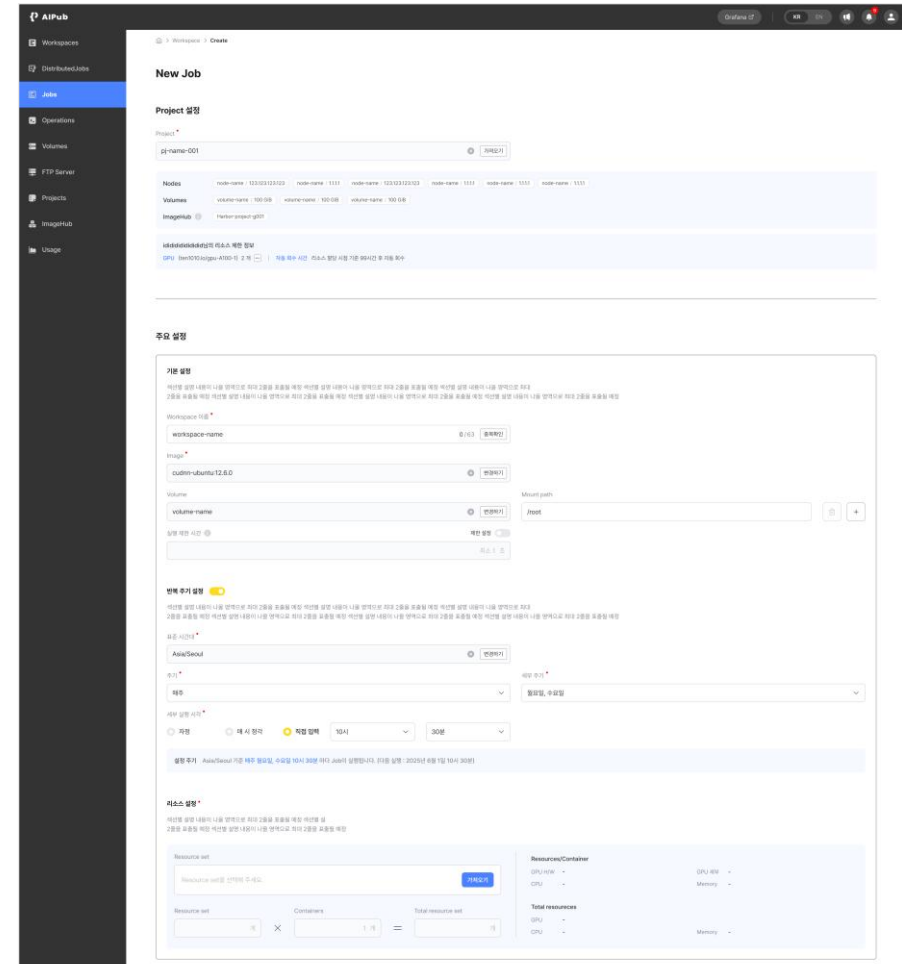
AI Pub Dev

AI Pub은 Kubernetes 기반이지만 직관적인 UI로 워크로드, 리소스, 볼륨을 통합 관리할 수 있어 비숙련자도 쉽게 접근 가능
실험 환경 구성, 실행, 결과 저장을 한 플랫폼 내에서 처리할 수 있어 전반적인 AI 연구 워크플로우가 단축됨

[Workspace 생성]



[Job 생성]



AI Pub은 Kubernetes 기반이지만 직관적인 UI로 워크로드, 리소스, 볼륨을 통합 관리할 수 있어 비숙련자도 쉽게 접근 가능
실험 환경 구성, 실행, 결과 저장을 한 플랫폼 내에서 처리할 수 있어 전반적인 AI 연구 워크플로우가 단축됨

[Workspace 상서]

AIPub

Admin

GrafanaKibana

> Workload > Workspaces

Workspaces 12

NewDelete

검색어를 입력해주세요

No	Name	Project	Owner	Node	SSH connection	Containers	Status	Age
1	test-edit-1	proj1	fe admin1 feadmin1	-	101.202.0.12:59409	0 / 3	Unavailable	18h 24m
2	request-to-vllm-2	proj1	정 수현 soohyun	pnode4.idc1.ten1010.io	101.202.0.12:63092	1 / 1	Available	21h 1m
3	workspacetest-3	proj1	qa admin qaadmin	vnode1.pnode8.idc1.ten1010.io	101.202.0.12:38542	0 / 1	Unavailable	23h 36m
4	workspace-test2	proj1	qa admin qaadmin	pnode4.idc1.ten1010.io	101.202.0.12:57504	1 / 1	Available	1d 3m
5	test-ssh-001	proj1	po admin1 podadmin1	pnode4.idc1.ten1010.io	101.202.0.12:33983	0 / 1	Unavailable	1d 27m
6	yunju-ws	proj1	정 윤주 podmin2	vnode1.pnode8.idc1.ten1010.io	101.202.0.12:11	3 / 3	Available	1d 16h
7	test-0625-1	proj1	fe admin1 feadmin1	-	101.202.0.12:30309	0 / 1	Unavailable	1d 18h
8	key-name-is-yunju	proj1	정 윤주 podmin2	vnode1.pnode8.idc1.ten1010.io	101.202.0.12:33882	3 / 3	Available	1d 22h
9	hjw	proj1	pm admin2 pmdmin2	pnode4.idc1.ten1010.io	101.202.0.12:17711	1 / 1	Available	2d 1h
10	yunju-available	proj1	정 윤주 podmin2	vnode1.pnode8.idc1.ten1010.io	101.202.0.12:13320	1 / 1	Available	2d 12h
11	yunju-workspace-2	proj1	정 윤주 podmin2	vnode1.pnode8.idc1.ten1010.io	101.202.0.12:40667	1 / 1	Available	2d 12h
12	example-workspace2	proj1	-	-	101.202.0.12:40033	0 / 2	Unavailable	2d 17h

50

Items per page

< 1 / 1 >

AIPub Admin
Grafana Kibana

[Workspace](#) > [Workspaces](#) > [Details](#)

workspace

[Edit](#)
[Pause](#)
[Delete](#)

Age
Available

4분 38초

(생성 일시 : 2025/06/27 11:50:46)

Image	harbor.cluster3.idc1.ten1010.io/common/vlmv0.7.0		
Environment	MY = hello		
Volume	Name	yunju-volume	
	Mount path	/root	
Address	Port	12345	
	Host	https://workspaceproj1.workspace.cluster4.idc1.ten1010.io/name	

Condition	Reason
-	-

Owner 윤주 정
ID poadmin2
Event

Project [proj1]

Nodes pnode4.idc1.ten1010.io vnode1.pnode6.idc1.ten1010.io vnode1.pnode8.idc1.ten1010.io

Resource

[129mcore-129mb-gpu-block-64mb-rtx-a2000] x 3

GPU ten1010.io/gpu-block-64mb-rtx-a2000 x 3
Type GPU VRAM 64MB

CPU 387m Core

Memory 387 MB

Containers

Available 3	Current 3	Desired 3
-------------	-----------	-----------

SSH connection Host 101.202.012 / Port 60388 / User root / CLI command ssh -i [your pem file path] -p 60388 root@101.202.012

[Public key 확인](#)

Containers 3

Read me

No	Name	Node	Status	Restart	Age	Action
1	workspace-a08350de-1	vnode1.pnode6.idc1.ten1010.io	Ready	0	4m 35s	Log Event
2	workspace-a08350de-0	vnode1.pnode6.idc1.ten1010.io	Ready	0	4m 36s	Log Event
3	workspace-a08350de-2	vnode1.pnode6.idc1.ten1010.io	Ready	0	4m 36s	Log Event

50 Items per page
< 1 / 1 >



AI Pub Ops



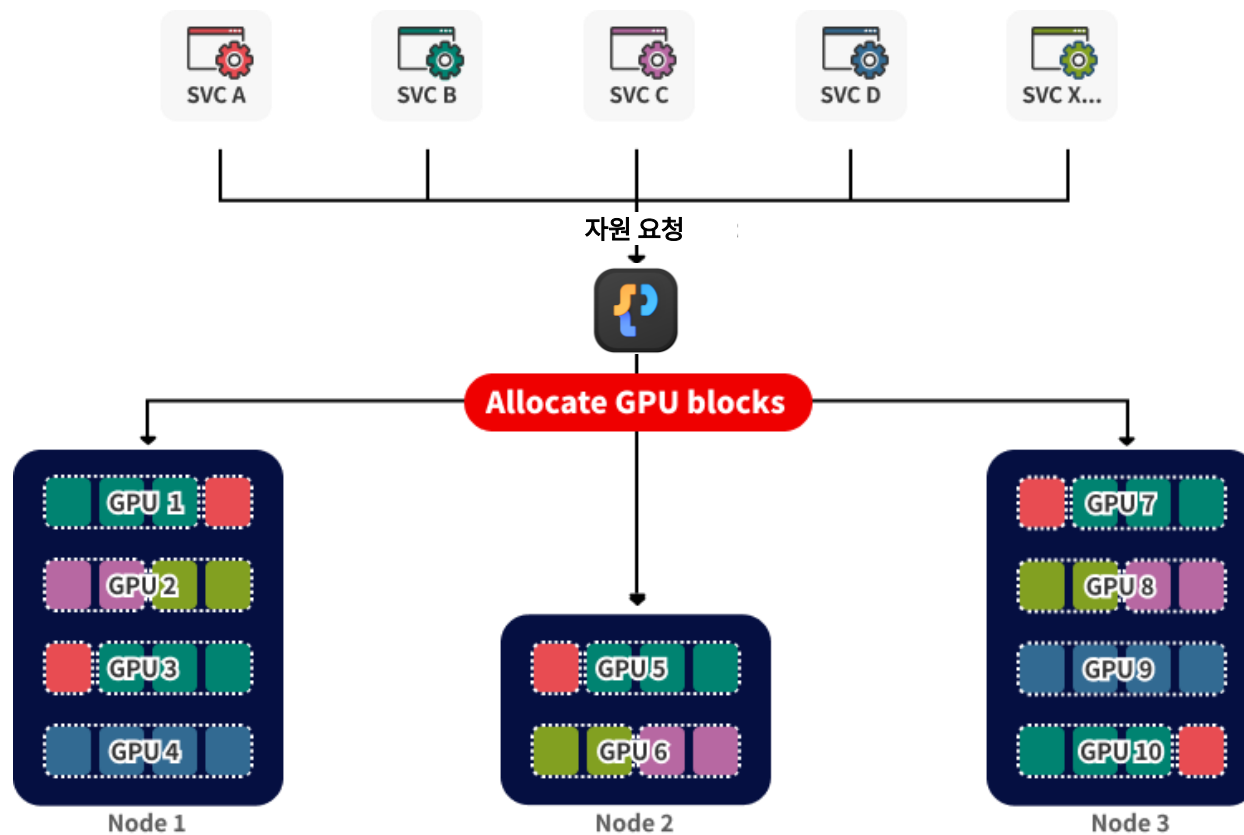
COASTER를 코어로 하여 AI 운영을 지원하는 완전 관리형 서비스

AlPub Ops 주요 기능 1.

사용자의 서비스를 이미지의 형태로 관리합니다.

GPU 1개를 100분할 하여 최소 단위로 운영해 비용 절감에 기여합니다.

비개발자도 서비스를 생성하거나 중지, 삭제할 수 있으며, 업데이트 및 롤백도 가능합니다.



AIPub Ops 주요 기능 2.

퍼블릭 클라우드의 **서비스 운영 및 관리 기능**을 제공하여
온프레미스 서버에서도 AI 서비스 운영을 돕는 기능을 지원합니다.

High Availability

이중화를 통해
Single point of failure 제거

Rolling Update

서비스 중단 없이
상시 업데이트 가능

Load Balancing

바쁘지 않은 서버로
서비스 요청을 자동 분배

Scale-out

서비스 요청에 따라
서버 수를 자동으로 늘려 트래픽 처리

Fail over 대응

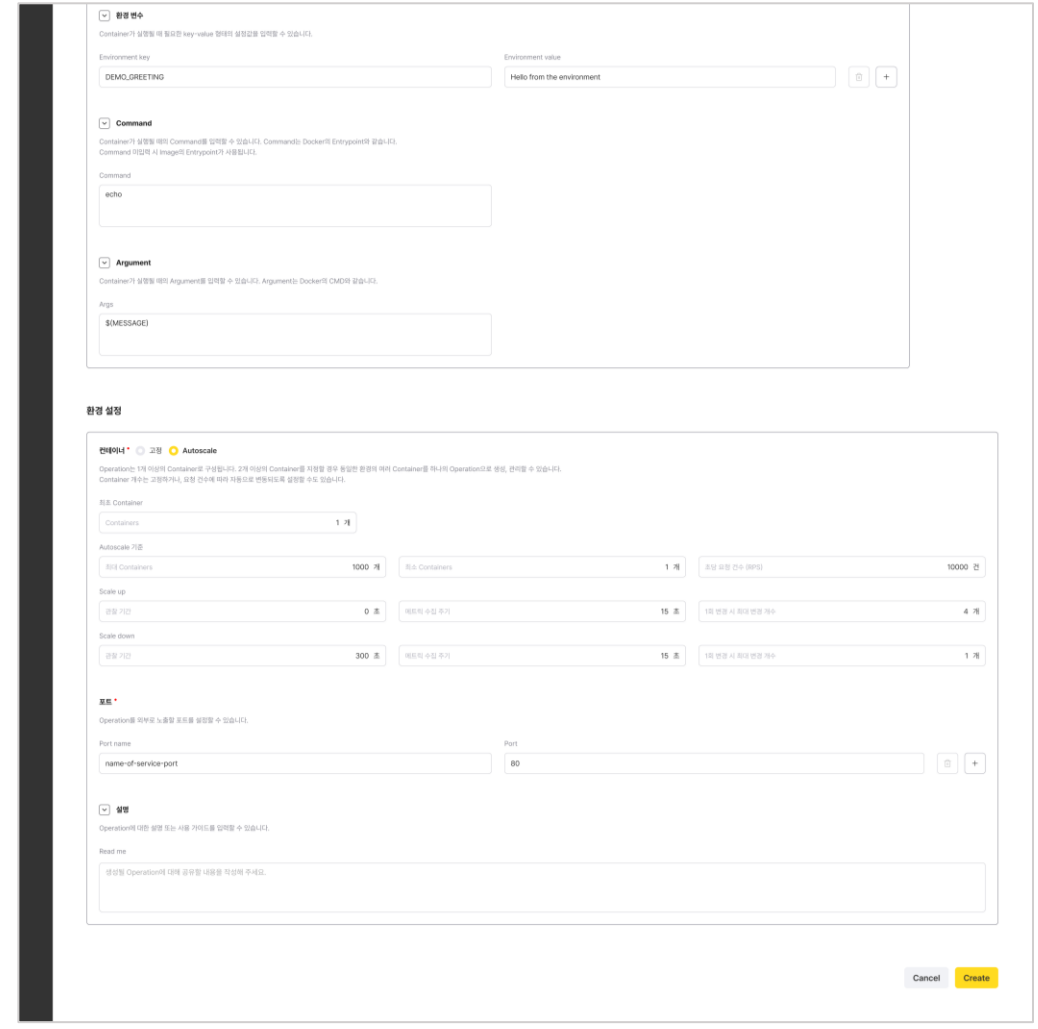
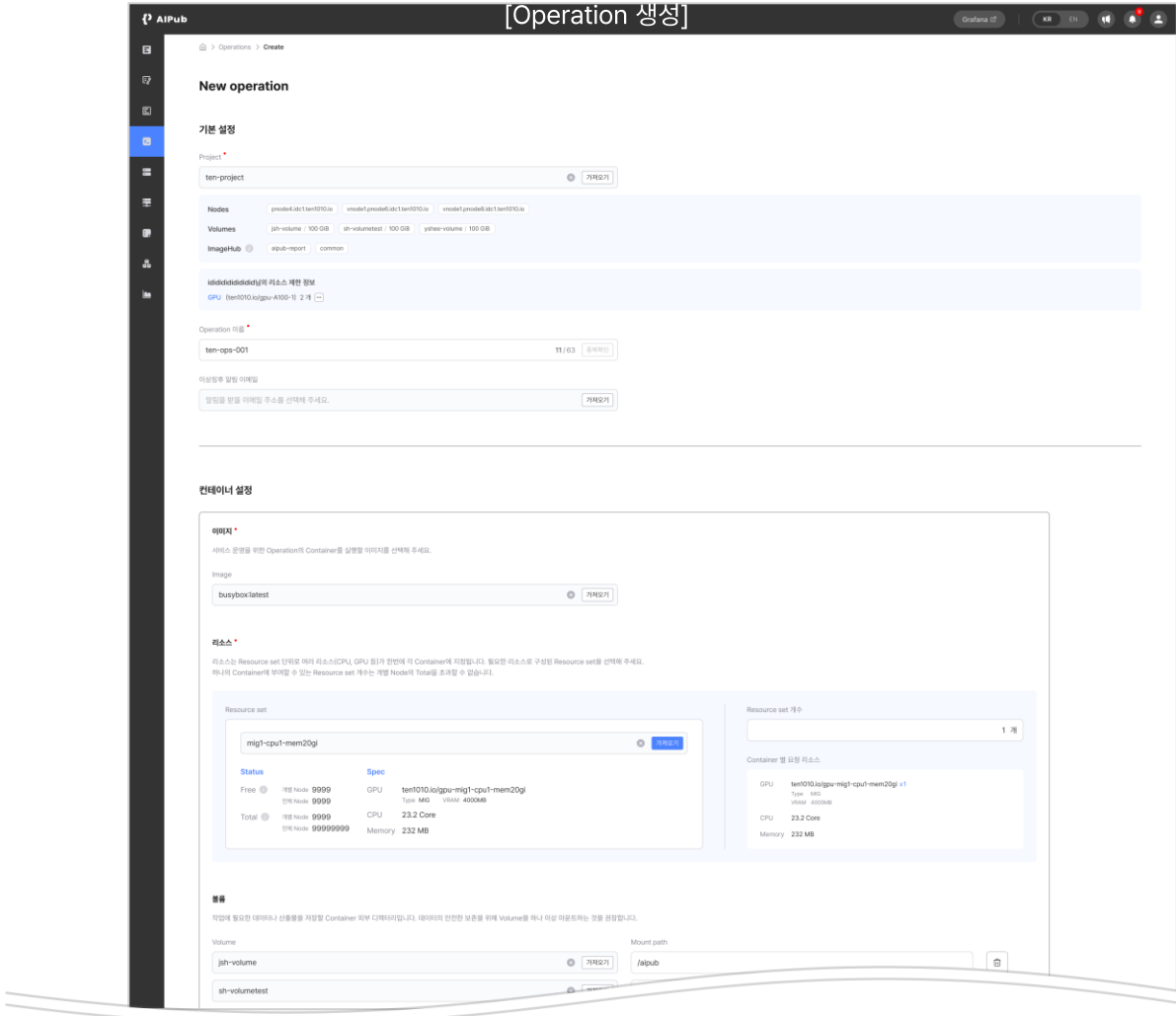
서비스 Fail over 시 탐지 기능 및
서비스를 새로 띄워 안정성 확보

이상징후 알림

중요 운영 이벤트 발생 시
카톡, 슬랙으로 실시간 알림

AI Pub Ops

AI Pub은 Auto-scaling 서빙과 예약 실행 기반 배치 작업을 UI로 쉽게 구성할 수 있어, AI 모델 운영과 지속적인 개선을 효율적이고 안정적으로 지원합니다.



AI Pub Ops

AI Pub은 Auto-scaling 서빙과 예약 실행 기반 배치 작업을 UI로 쉽게 구성할 수 있어, AI 모델 운영과 지속적인 개선을 효율적이고 안정적으로 지원합니다.

[Operation 상세]

Kubernetes
Workspaces DistributedJobs Jobs

Operations

Volumes FTP Server Projects ImageHub Usage

> Workload > Operations > Details
Grafana ⌵ K8S ⌵ 🔔

Namespace: devteam
{ops-name}

Edit
Pause
Delete

Age

Available

105일 13시간 35분 16초

(상장 일시 : 2019/10/28 10:28:30)

Image	agor.io/linkerd-io/controller:16.04
Environment	key = value
Command	/bin/bash
Argument	/bin/bash
Address	Port: 8090 ☐ Host: a.b.c.d.com:12345/tpl-port1
Volume	Name: volume-yuntje-1 ☐ Mount path: root/agor.io/linkerd-io/control

Condition Reason -

Message -

Owner: kim tan
ID: kintan010
Event
Monitoring

Project ⌵

[ten-project]

Nodes: nd-001 nd-002 nd-003 nd-003

Resource

[m1g1-cpu1-mem20gi] x 9

CPU: ten010.klgw-m1g1-cpu1-mem20gi x 9
Type: m3 VMWare: acorn3e

CPU: 45.4 Core

Memory: 454 MB

Containers

Min / Max: 1 / 9 Autoscale 기제 (RPS): 30

Window: up 300s / down 300s Period: up 300s / down 300s At once: up 3 / down 1

알림대상
홍길동 / hong@email.com
홍길동 / hong@email.com
홍길동 / hong@email.com

Replicas 10

Versions 10

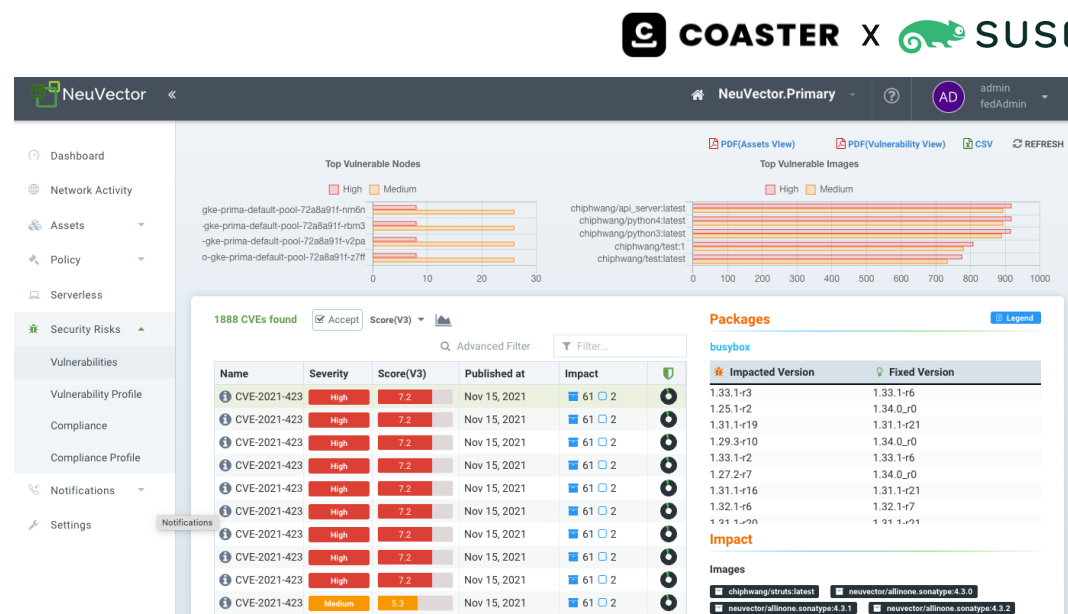
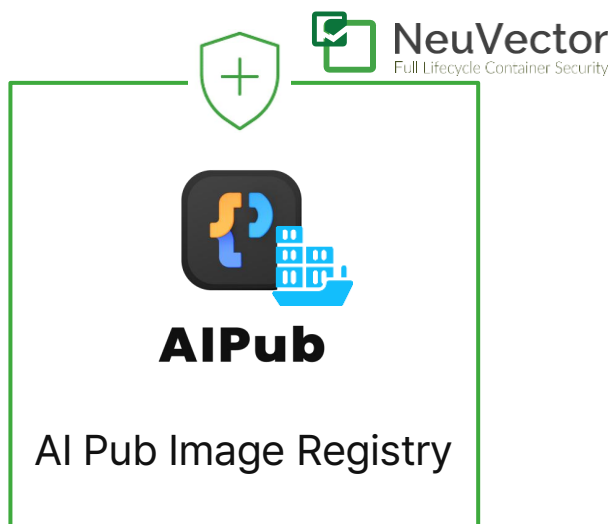
Read me

No	Version ⌵	Revision ⌵	Image	Container ⌵	Creation ⌵	Action
1	99	[revision name]	[image name]	99.999	2024/10/10 00:00:00	<button>Rollback ⌵</button>
2	98	test-operation-7c313818-9	agor.io/linkerd-io/controller:16.04	99.999	2024/10/10 00:00:00	<button>Rollback</button>
3	97	test-operation-7c313818-8	agor.io/linkerd-io/controller:16.04	99.999	2024/10/10 00:00:00	<button>Rollback</button>
4	96	test-operation-7c313818-7	agor.io/linkerd-io/controller:16.04	99.999	2024/10/10 00:00:00	<button>Rollback</button>
5	95	test-operation-7c313818-6	agor.io/linkerd-io/controller:16.04	99.999	2024/10/10 00:00:00	<button>Rollback</button>
6	94	test-operation-7c313818-5	agor.io/linkerd-io/controller:16.04	99.999	2024/10/10 00:00:00	<button>Rollback</button>
7	93	test-operation-7c313818-4	agor.io/linkerd-io/controller:16.04	99.999	2024/10/10 00:00:00	<button>Rollback</button>
8	92	test-operation-7c313818-3	agor.io/linkerd-io/controller:16.04	99.999	2024/10/10 00:00:00	<button>Rollback</button>
9	91	test-operation-7c313818-2	agor.io/linkerd-io/controller:16.04	99.999	2024/10/10 00:00:00	<button>Rollback</button>
10	90	test-operation-7c313818-1	agor.io/linkerd-io/controller:16.04	99.999	2024/10/10 00:00:00	<button>Rollback</button>

50 ▾
Items per page
< 1 / 1

AI Pub Ops는 NeuVector의 스캐닝 기능을 사용하여 다양한 경로로 유입되는 도커 이미지에 대한 보안 필터링이 가능합니다.

컨테이너 이미지, 클러스터를 구성하는 노드, 동작 중인 컨테이너의 취약점을 검사할 수 있습니다.



CVE를 기준으로 스캐닝

CVE는 다양한 취약점을 통일해 고유한 넘버링을 부여한 것으로 장치, 시스템, 프로그램 해킹에 악용될 수 있는 컴퓨터 보안 취약성 및 시스템 결함에 코드를 부여하고 관리



AI Pub K10s



K8s 숙련자를 위한 Kubernetes Native Platform

AI Pub K10s

AI Pub의 Kubernetes UI는 CLI만 사용하는 방식 대비 전체 리소스 상태를 한눈에 파악하고, 검색·관리·분석까지 직관적으로 수행할 수 있어 숙련자들의 운영 효율성을 극대화합니다.

[Pod 목록]

AI Pub

Admin

Workload

Pods

Deployments

ReplicaSets

StatefulSets

DaemonSets

Jobs

CronJobs

Operations

Network & Service

Services

Ingresses

IngressClasses

NetworkPolicies

Endpoints

Configuration

ConfigMaps

Secrets

HPAs

PodDisruptionBudgets

PriorityClasses

Storage

PVs

PVCs

StorageClasses

SFTPServers

Cluster Mgmt

Extra Cluster Mgmt

Users & Auth

Show Permissions

Pods

Namespace

+ Create

Delete

Search

No	Name	Namespace	Ready	Status	Restarts	Node	Age
1	workspace-a08350de-2	proj1	0/1	Pending	0	vnode1.pnode6.idc1...	25m 32s
2	workspace-a08350de-1	proj1	1/1	Running	0	vnode1.pnode6.idc1...	31m 33s
3	workspace-a08350de-0	proj1	1/1	Running	0	vnode1.pnode6.idc1...	31m 34s
4	pod-0	proj1	0/1	Running	16	pnode4.idc1.ten101...	57m 31s
5	testpod-vshee	proj1	0/1	Running	21	pnode4.idc1.ten101...	1h 26m
6	workspace-test-3-ecf99105-0	proj1	0/1	Pending	0	vnode1.pnode8.idc1...	21h 13m
7	test-ssh-5774r	proj1	1/1	Running	0	pnode4.idc1.ten101...	21h 25m
8	test-ssh-nvhtt	proj1	1/1	Running	0	pnode4.idc1.ten101...	21h 25m
9	request-to-vlim-2-d06f7143-0	proj1	1/1	Running	0	pnode4.idc1.ten101...	21h 42m
10	gpu-test-job-3-wf46n	proj1	0/1	Succeeded	0	pnode4.idc1.ten101...	23h 18m
11	workspace-test2-3aa12042-0	proj1	1/1	Running	0	pnode4.idc1.ten101...	1d 43m
12	gpu-test-job-2-zwnc7	proj1	0/1	Succeeded	0	pnode4.idc1.ten101...	1d 1h
13	gpu-test-job-tcbz4	proj1	0/1	Succeeded	0	pnode4.idc1.ten101...	1d 1h
14	test-ssh-001-2ecaddcb-0	proj1	0/1	Pending	0	pnode4.idc1.ten101...	1d 1h
15	key-name-is-yunju-9b5df860-1	proj1	1/1	Running	0	vnode1.pnode8.idc1...	1d 1h
16	key-name-is-yunju-9b5df860-2	proj1	1/1	Running	0	vnode1.pnode8.idc1...	1d 1h
17	key-name-is-yunju-9b5df860-0	proj1	1/1	Running	0	vnode1.pnode8.idc1...	1d 1h
18	yunju-ws-37907f17-1	proj1	1/1	Running	0	vnode1.pnode8.idc1...	1d 17h
19	yunju-ws-37907f17-2	proj1	1/1	Running	0	vnode1.pnode8.idc1...	1d 17h
20	yunju-ws-37907f17-0	proj1	1/1	Running	0	vnode1.pnode8.idc1...	1d 17h

0 of 28 row(s) selected.

Rows per page 50

Page 1 of 1

<< < > >>

AI Pub K10s

AI Pub의 Kubernetes UI는 CLI만 사용하는 방식 대비 전체 리소스 상태를 한눈에 파악하고, 검색·관리·분석까지 직관적으로 수행할 수 있어 숙련자들의 운영 효율성을 극대화합니다.

[Yaml Editor]

The image displays two side-by-side screenshots of the AI Pub Kubernetes UI, specifically the 'Create Pod' and 'Edit Pod' interfaces. Both interfaces feature a sidebar on the left with a navigation menu including 'Workload', 'Network & Service', 'Configuration', 'Storage', 'Cluster Mgmt', 'Extra Cluster Mgmt', 'Users & Auth', and 'Show Permissions'. The main area contains a 'Create Pod' or 'Edit Pod' form with a 'Namespace' dropdown and a 'Yaml Editor' section. The 'Create Pod' interface shows a YAML editor with fields for 'apiVersion', 'kind', 'metadata', 'spec', and 'containers'. The 'Edit Pod' interface shows a similar editor with a 'Hide Diff' button and a 'Yaml Editor' section. The 'Edit Pod' interface also includes a 'Resources' section with a 'Show Permissions' button. The 'Edit Pod' interface shows a diff view of the YAML configuration, with changes highlighted in green and red. The 'Edit Pod' interface also includes a 'Resources' section with a 'Show Permissions' button. The 'Edit Pod' interface shows a diff view of the YAML configuration, with changes highlighted in green and red. The 'Edit Pod' interface also includes a 'Resources' section with a 'Show Permissions' button.

AI Pub K10s

다양한 사용자 환경(Web UI, YAML editor, CLI)에서 동일하게 생성 및 편집 가능하며 이를 하나의 GUI에서 관리

[Operation 예시: YAML editor]

[Operation 예시: Form UI]

```
lyj@vnode5:~/lyj$ cat ten-ai-serving.yaml
apiVersion: aipub.ten1010.io/v1alpha1
kind: Operation
metadata:
  name: ten-ai-serving
  namespace: aipub
  annotations:
    resourceType.aipub.ten1010.io/request: cpu-10m-memory-128mi=1
spec:
  replicas: 1
  ports:
    - name: port
      port: 8090
  autoScaling:
    maxReplicas: 10
    minReplicas: 1
    rps: "1"
  pod:
    image: "python:3.12-slim"
    command:
      - "/bin/sh"
      - "-c"
      - "sleep infinity"

lyj@vnode5:~/lyj$ sudo kubectl apply -f ten-ai-serving.yaml
operation.aipub.ten1010.io/ten-ai-serving created
lyj@vnode5:~/lyj$ k get ops -n aipub
NAME                                AVAILABLE  ADDRESSES                                                                 REPLICAS  AVAILABLE_REPLICAS
operation-yunju                     false     ["https://controller.operation.cluster9.idcl.ten1010.io/aipub-operation-yunju-portname"]  2
ten-ai-serving                      true      ["https://controller.operation.cluster9.idcl.ten1010.io/aipub-ten-ai-serving-port"]      1          1
test-test-ssh2                     false     ["https://controller.operation.cluster9.idcl.ten1010.io/aipub-test-test-ssh2-test"]      1
```

[Operation 예시: CLI]



RA:X

AI Reference Architecture of TEN



TEN의 노하우를 더한 테스트 기반의 AI 인프라 구축 컨설팅 서비스

하드웨어를 예산에 맞춰 각각 구입하도록 제안하지 않습니다.
인프라 **전체**의 **가치**를 볼 수 있게 하겠습니다.

“저희는 Llama 4를
학습할 계획입니다.”



고객

모델, 데이터 샘플

AI Pub &
AI 인프라 구성 제안



성능 측정치

측정치에 맞는
HW 견적

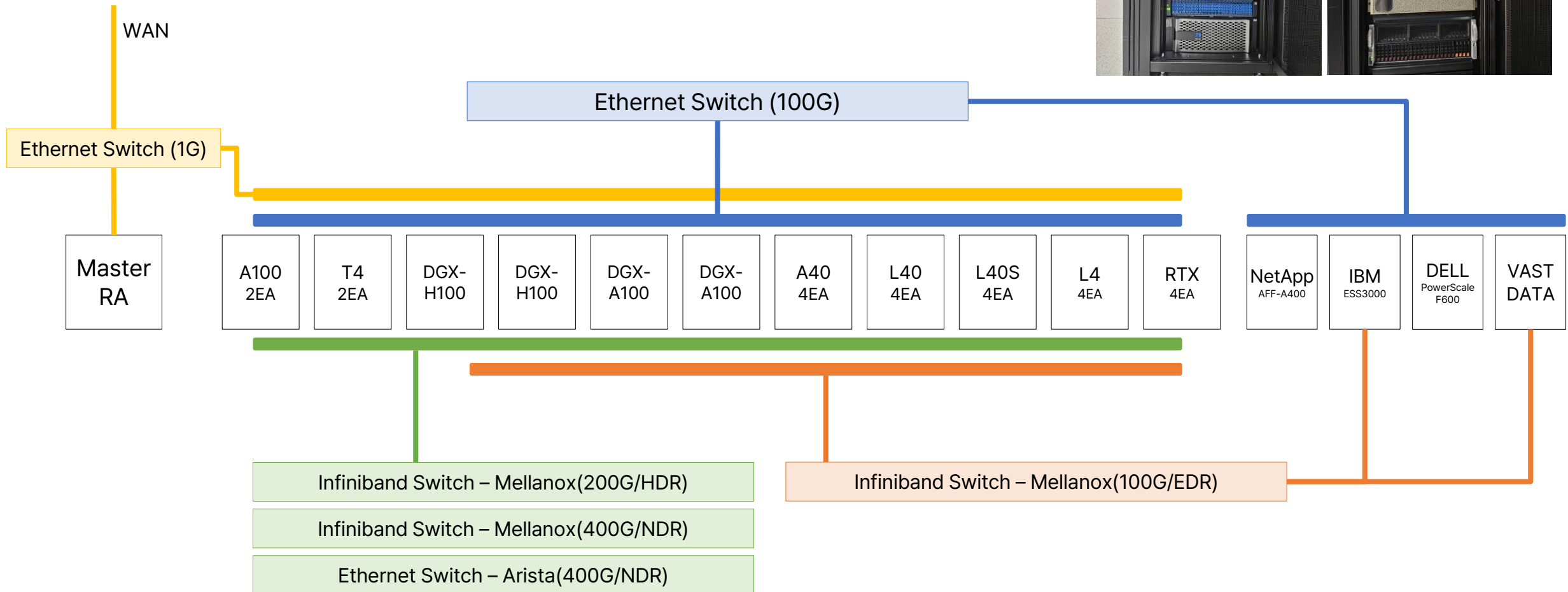
“그 정도 수치면 ”
이 HW 가 알맞습니다.



하드웨어 벤더

TEN이 직접 구축하고 운영하는 Reference Architecture을 소개합니다.

최신 GPU 서버와, 스토리지, 네트워크로 구성된 레퍼런스 아키텍처를 구성하였습니다.
AI 모델 별 학습과 운영을 위한 컴퓨팅 리소스 측정이 가능합니다.



RA:X 를 통해 DeepSeek의 서비스 운영 인프라 구성을 비용 효율적으로 따져보기

DeepSeek-V3, R1

파라미터 개수: 671B

활성화 파라미터: 37B

**8 NVIDIA H200 GPUs providing
1128 GB of GPU memory**

Model	AIME 2024 pass@1	AIME 2024 cons@64	MATH-500 pass@1	GPQA Diamond pass@1	LiveCodeBench pass@1	CodeForces rating
GPT-4o-0513	9.3	13.4	74.6	49.9	32.9	759
Claude-3.5-Sonnet-1022	16.0	26.7	78.3	65.0	38.9	717
o1-mini	63.6	80.0	90.0	60.0	53.8	1820
QwQ-32B-Preview	44.0	60.0	90.6	54.5	41.9	1316
DeepSeek-R1-Distill-Qwen-1.5B	28.9	52.7	83.9	33.8	16.9	954
DeepSeek-R1-Distill-Qwen-7B	55.5	83.3	92.8	49.1	37.6	1189
DeepSeek-R1-Distill-Qwen-14B	69.7	80.0	93.9	59.1	53.1	1481
DeepSeek-R1-Distill-Qwen-32B	72.6	83.3	94.3	62.1	57.2	1691
DeepSeek-R1-Distill-Llama-8B	50.4	80.0	89.1	49.0	39.6	1205
DeepSeek-R1-Distill-Llama-70B	70.0	86.7	94.5	65.2	57.5	1633

Inter-Token latency simulation

10ms

LLM 서비스를 제공하기 위해서는 초고속 연산 자원과 대용량 데이터 처리 능력이 필수적이어서 막대한 비용이 소요된다. 이로 인해 각 모델의 특성과 운영 환경에 맞는 최적의 인프라 구성이 절실하며, 효율성과 안정성을 동시에 갖춘 시스템 구축이 요구된다.

주식회사 Ten에서는 RAX 서비스를 통해 고객이 사용하려는 모델에 최적화된 효율적인 인프라를 구축할 수 있도록 지원하며, 비용 절감과 운영 효율성 극대화를 도모한다. 당사의 솔루션은 주어진 예산 내에서 최적의 성능을 발휘하여 LLM 서비스의 경쟁력을 강화하고, 고객의 비즈니스 성공을 견인하는 핵심 전략으로 자리매김할 것이다.

50ms



100ms

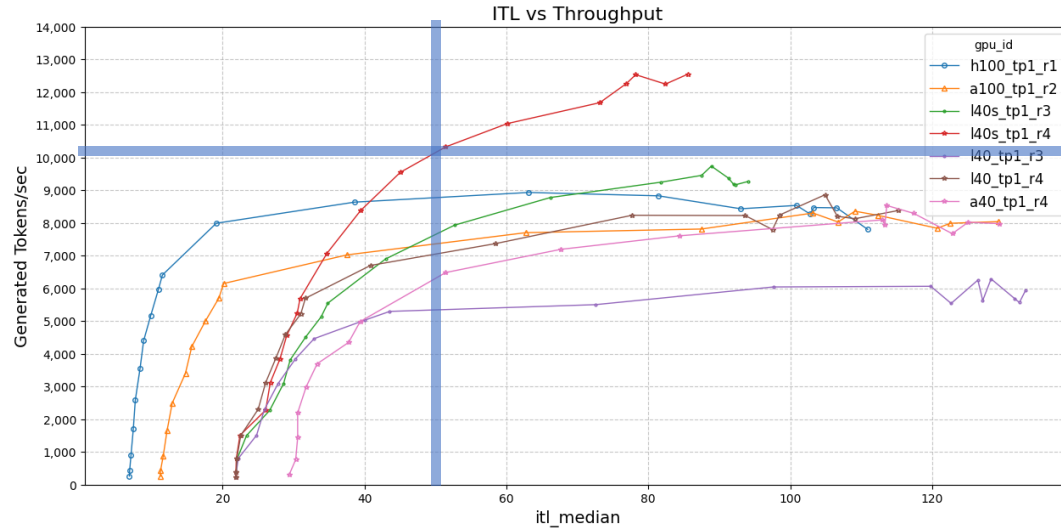


250ms

LLM 서비스를 제공하기 위해서는 초고속 연산 자원과 대용량 데이터 처리 능력이 필수적이어서 막대한 비용이 소요된다. 이로 인해 각 모델의 특성과 운영 환경에 맞는 최적의 인프라 구성이 절실하며, 효율성과 안정성을 동시에 갖춘 시스템 구축이 요구된다.

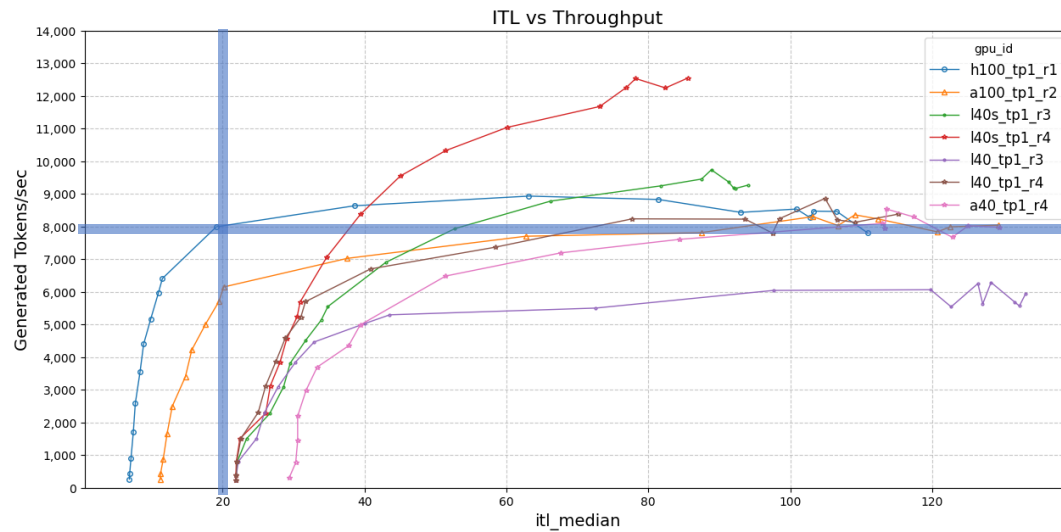
주식회사 Ten에서는 RAX 서비스를 통해 고객이 사용하려는 모델에 최적화된 효율적인 인프라를 구축할 수 있도록 지원하며, 비용 절감과 운영 효율성 극대화를 도모한다. 당사의 솔루션은 주어진 예산 내에서 최적의 성능을 발휘하여 LLM 서비스의 경쟁력을 강화하고, 고객의 비즈니스 성공을 견인하는 핵심 전략으로 자리매김할 것이다.

DeepSeek-R1-Distill-Qwen-7B 모델의 최적의 운영 GPU 자원 제안



- Target Inter-Token Latency = 50ms

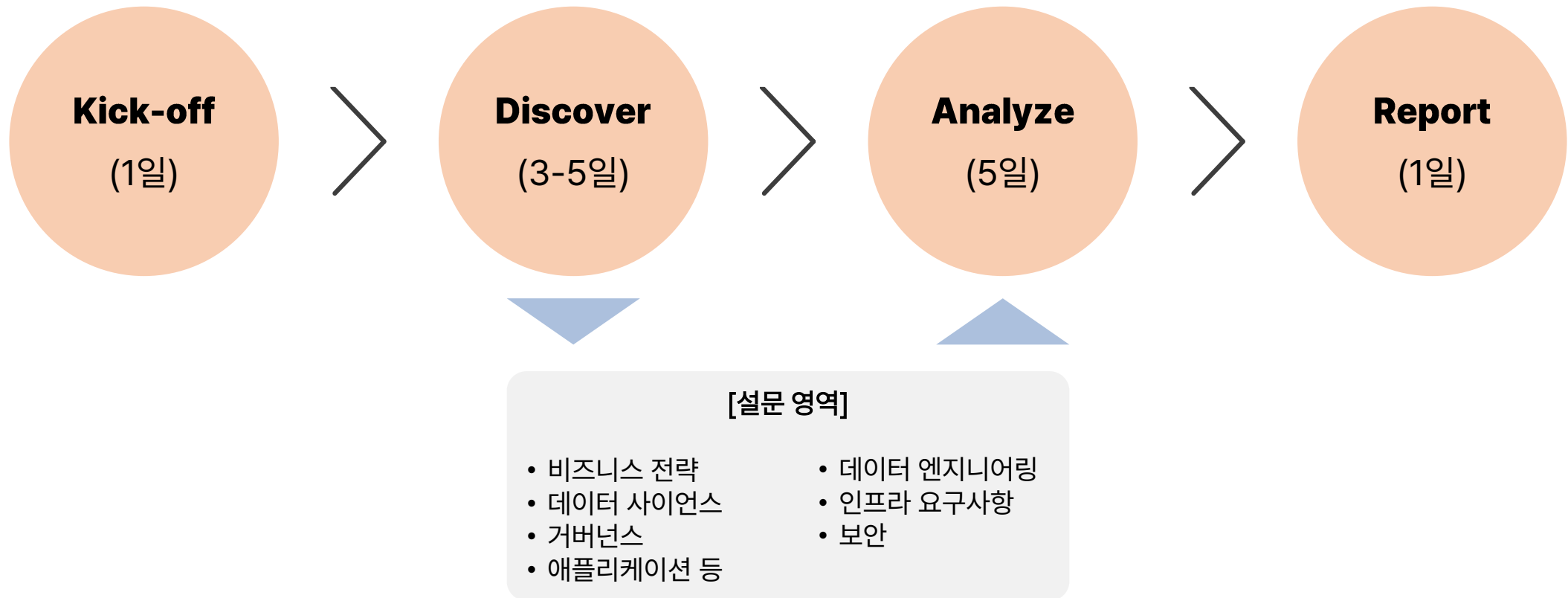
→ L40s 4개



- Target Inter-Token Latency = 20ms

→ H100 1개

RA:X 를 기반으로 한 인프라 구성 컨설팅은 아래의 프로세스로 진행됩니다.



TEN은 AI 인프라 구성의 어려움과 해결 방안에 대한 세미나, 워크숍을 개최합니다.
인프라에 대한 고민에 공감하고, 해결을 돕기 위한 **TEN**의 목소리를 다양한 곳에 전합니다.

**AI 인프라와
NVIDIA 인증 솔루션
통합 구성 워크숍**

나에게 딱 맞는 AI 인프라 조합을 찾는 여정



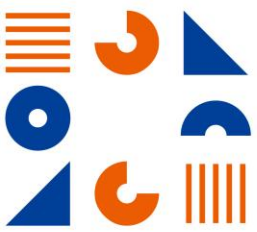
RA:X

Nov. 23 Thu 14:00~17:00

드림플러스 강남

**AI 인프라와
NVIDIA 인증 솔루션
통합 구성 워크숍**

나에게 딱 맞는 AI 인프라 조합을 찾는 여정



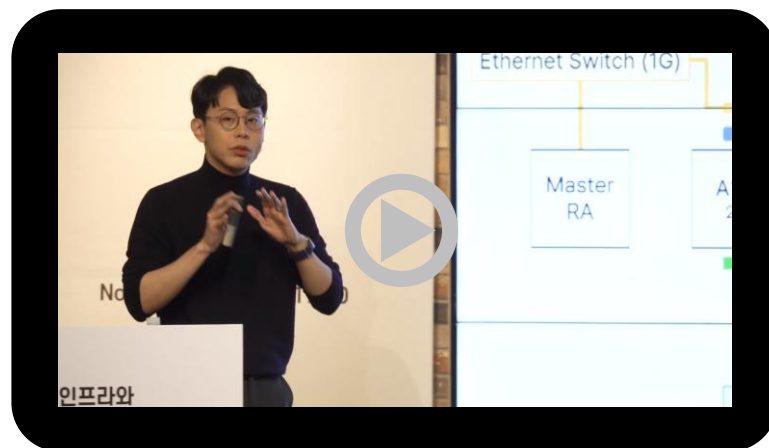
RA:X

Nov. 23 Thu 14:00~17:00

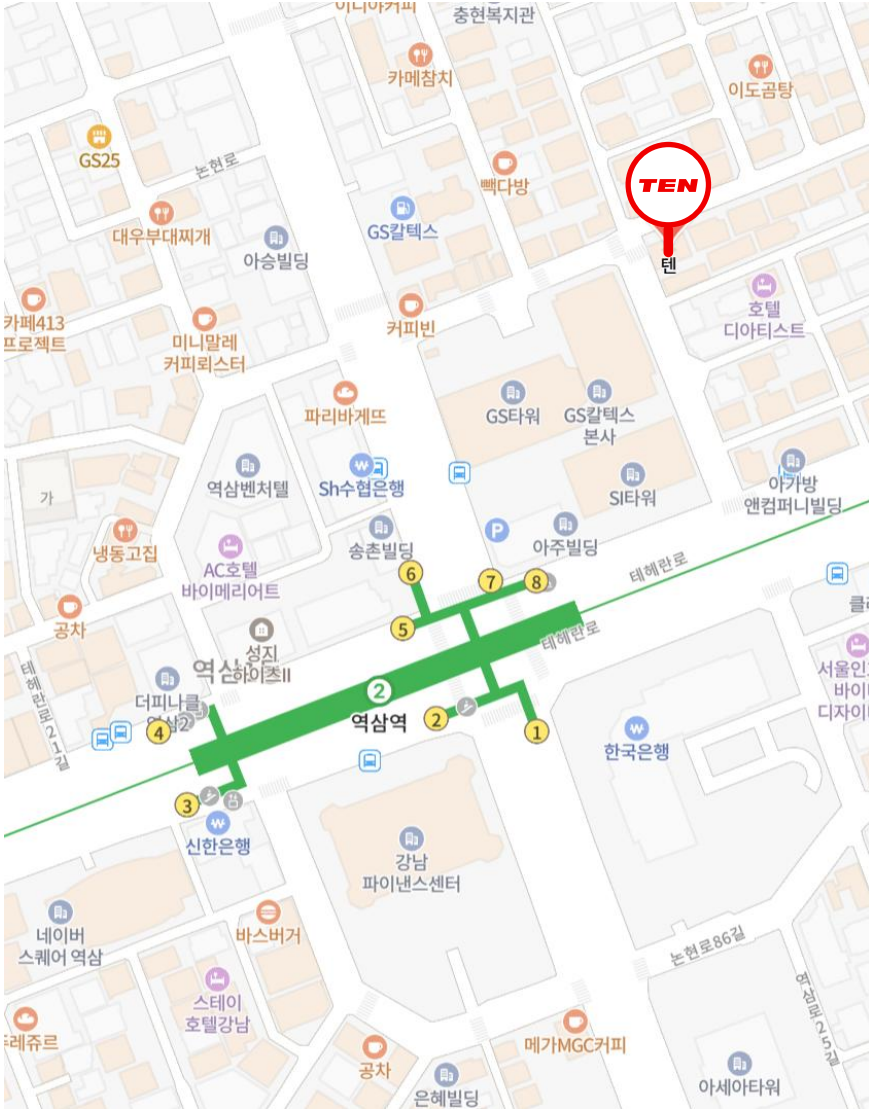
드림플러스 강남



2023.11
AI 인프라와 NVIDIA 인증 솔루션
통합 구성 워크숍
[Opening Session by TEN](#)



2023.11
AI 인프라와 NVIDIA 인증 솔루션
통합 구성 워크숍
[Session C, Closing Session by TEN](#)



**MORSE ME
WHEN YOU NEED
HELP WITH AI.**

경쟁력 있는 AI를 위한 파트너가 필요하신가요?

고효율 솔루션과 신뢰도 높은 파트너십을 갖춘 주식회사 텐과의 협업으로
이전과 다른 AI 개발·운영을 경험하실 수 있습니다.

A 서울특별시 강남구 테헤란로27길 18, 5층

P +82-2-6956-1071

M helloten@ten1010.io

H <https://ten1010.io/>



**AI로운
세상을 향해.**

THANK YOU

Copyright © 2024 Ten Inc. All rights reserved