

You Focus on Your AI Missions
We Handle the AI Environment

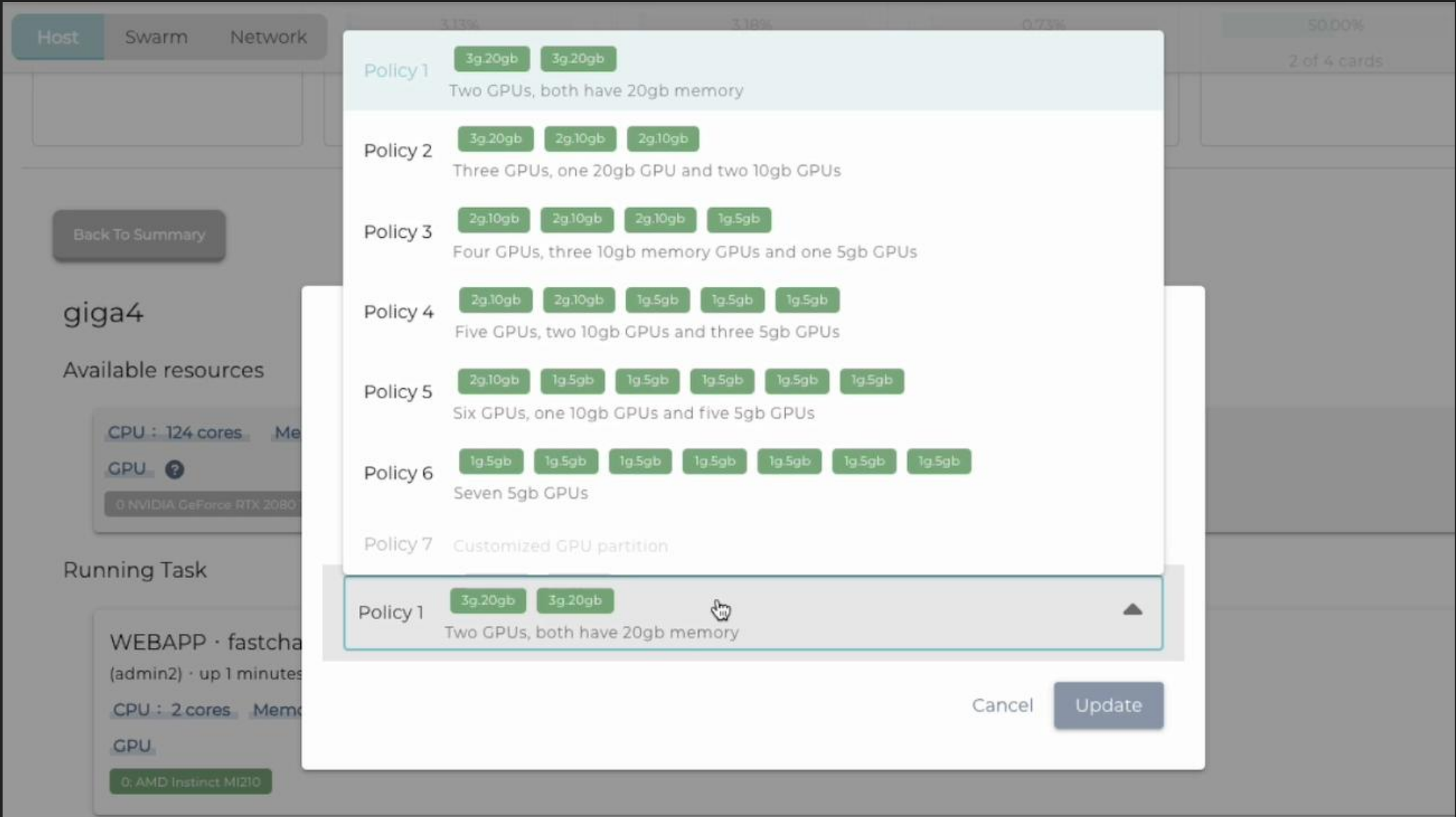
從硬體到模型
一站到底的AI行動方案

BDM Brad Huang
業務發展經理 黃丞毅

前言



前言



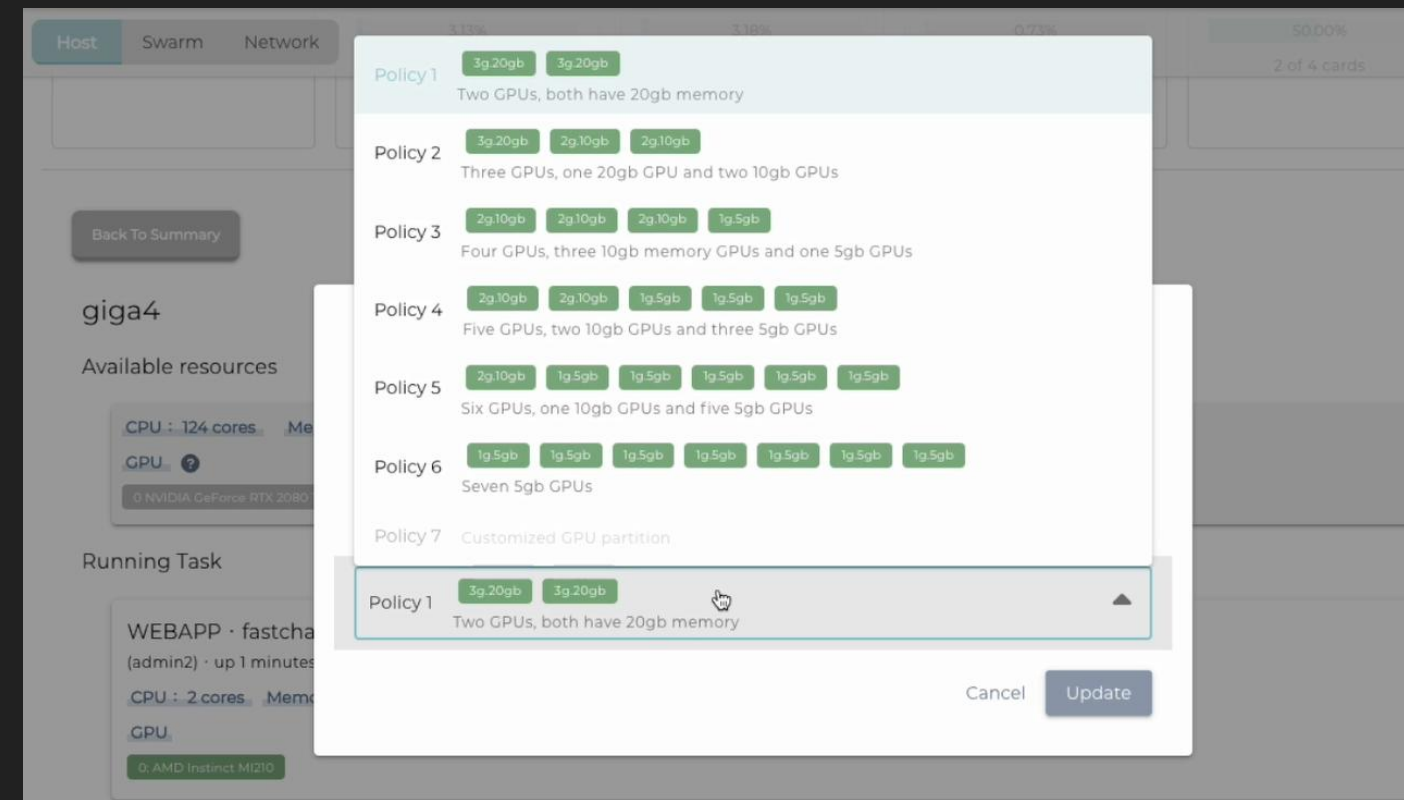
前言



前言



時效性



操作直觀性



整合性

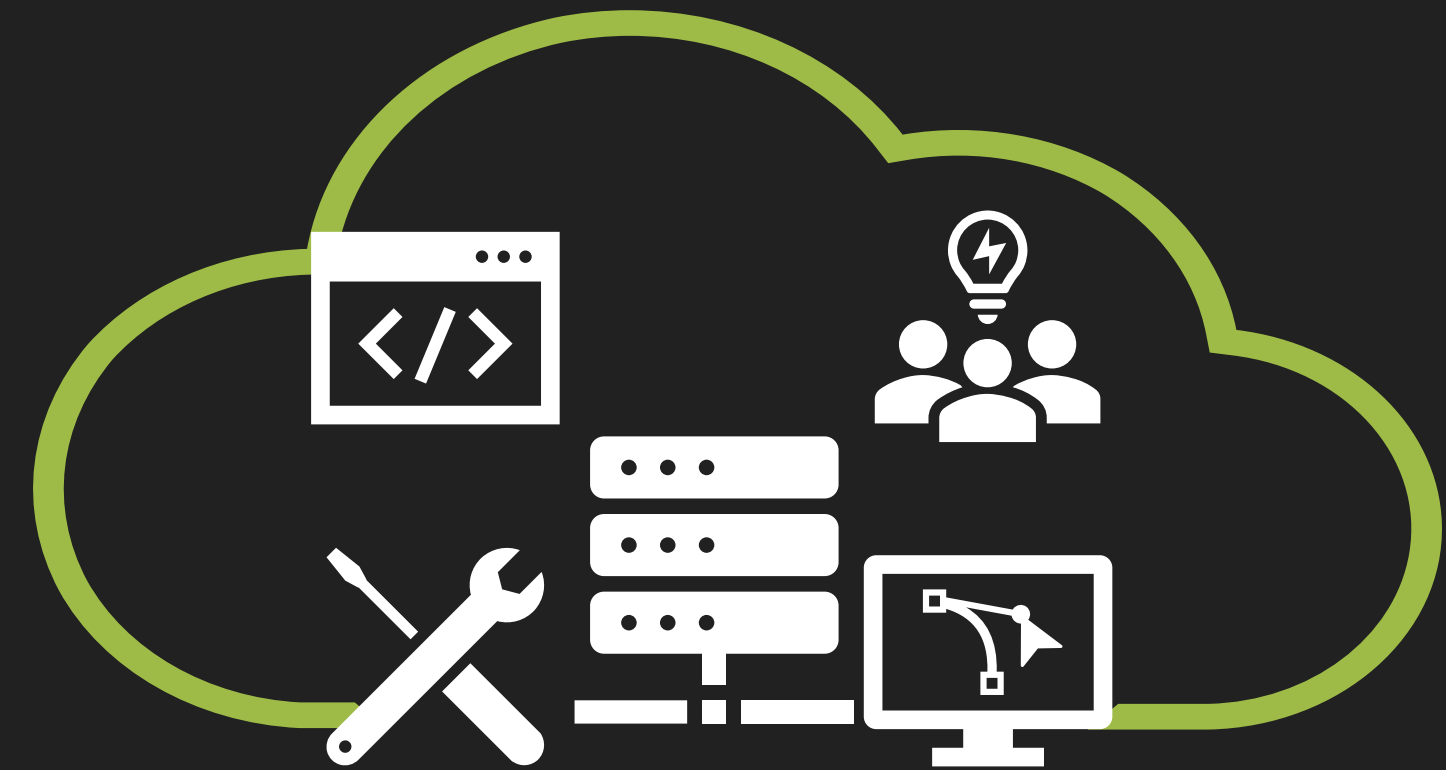
開發者的運算服務需求

Bare Metal as a Service
(or Infrastructure as a Service / IaaS)



以提供硬體算力
服務為主

AI as a Service
(or Software as a Service / SaaS)

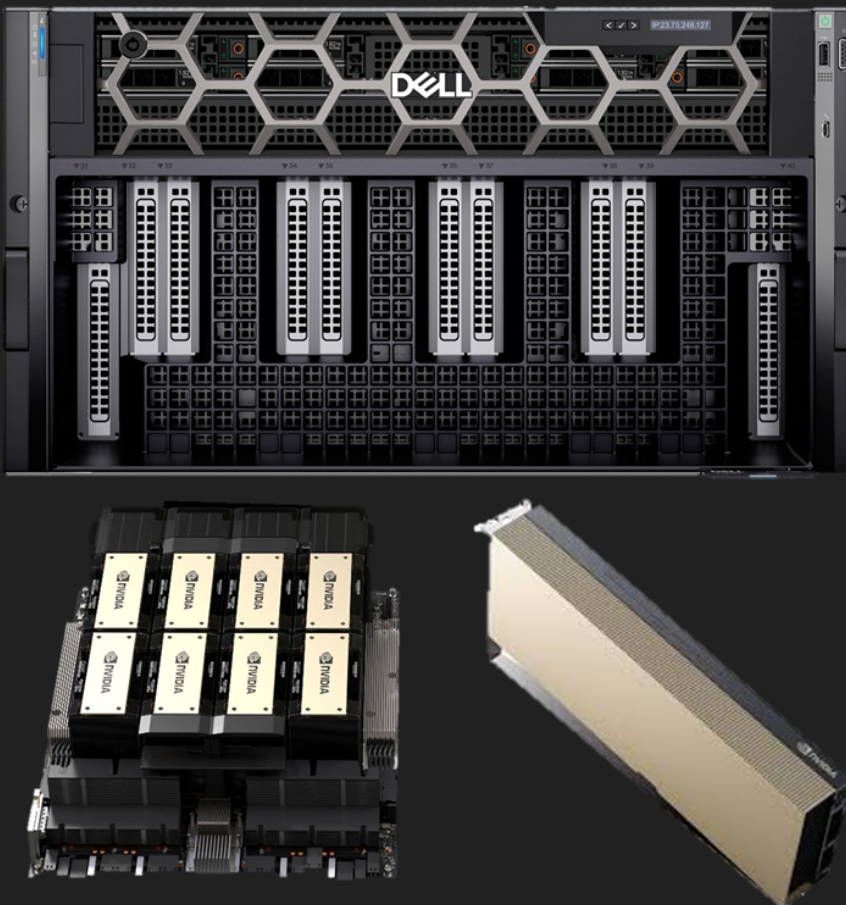


- 硬體算力服務
- AI 開發IDE與各式工具
- AI 模型樣板
- AI 開發環境
- 團隊協同開發合作

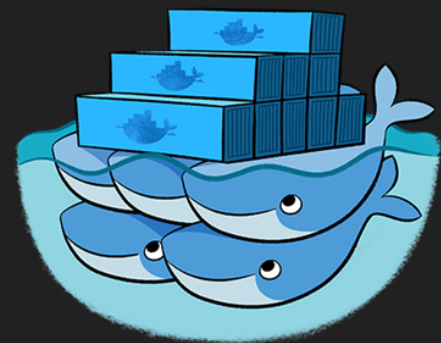
管理節點



運算節點



容器化管理



AI / MLOps 開發平台

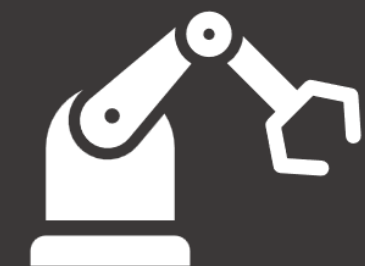
管理



整合



介接



從硬體資源到軟體整合
從模型開發到推論服務
一站式・完整開發環境

儲存節點



組織既有

- 帳號管理系統
- 計費管理系統
- 通訊管理系統

AI開發平台需求對象

IT / MIS團隊

- ✓ 系統叢集與備援
- ✓ GPU算力分割
- ✓ 硬體資源限制
- ✓ 訓練排程規劃
- ✓ 開發者權限管控

模型開發團隊

- ✓ 自動標註功能
- ✓ 開發IDE與標註工具整合於平台
- ✓ 支援NVAIE / NIM與開源模型樣板及框架
- ✓ 資料集與模型版本管控
- ✓ 開發流程標準化

AI維運團隊

- ✓ 可於平台直接啟動模型推論服務
- ✓ 模型加密機制
- ✓ 自動啟動模型推論服務功能
- ✓ 模型效能監控

平台在模型開發過程所支援的各項功能

平台管理

- 平台將管理、運算、儲存節點叢集納管，並進行監控、分配，支援跨節點運算
- 可透過平台GUI進行GPU分割
- 平台可自主設定運算資源限制，亦可限制GPU使用數量及時間
- 平台支援開發者權限、資料集、模型版本管控，設定使用排程並記錄使用時間



資料處理

- 完整資料處理功能支援資料視覺化預覽
- 透過介面將桌面資料拖曳上傳至平台

資料標註

- 內建常用資料標註工具
- 資料自動標註功能

模型佈署

- 啟用WebApp功能於平台以API、IP方式介接前端系統，或將模型打包成image
- 模型加密機制，提升模型安全性

流程標準化

- 使用平台Pipeline功能建構模型產出標準化流程，以利模型優化與模型再利用

模型評估

- 平台介面訓練數值可視化並可比較多筆實驗結果
- 將驗證資料集透過平台餵進程式碼產生驗證結果

模型訓練

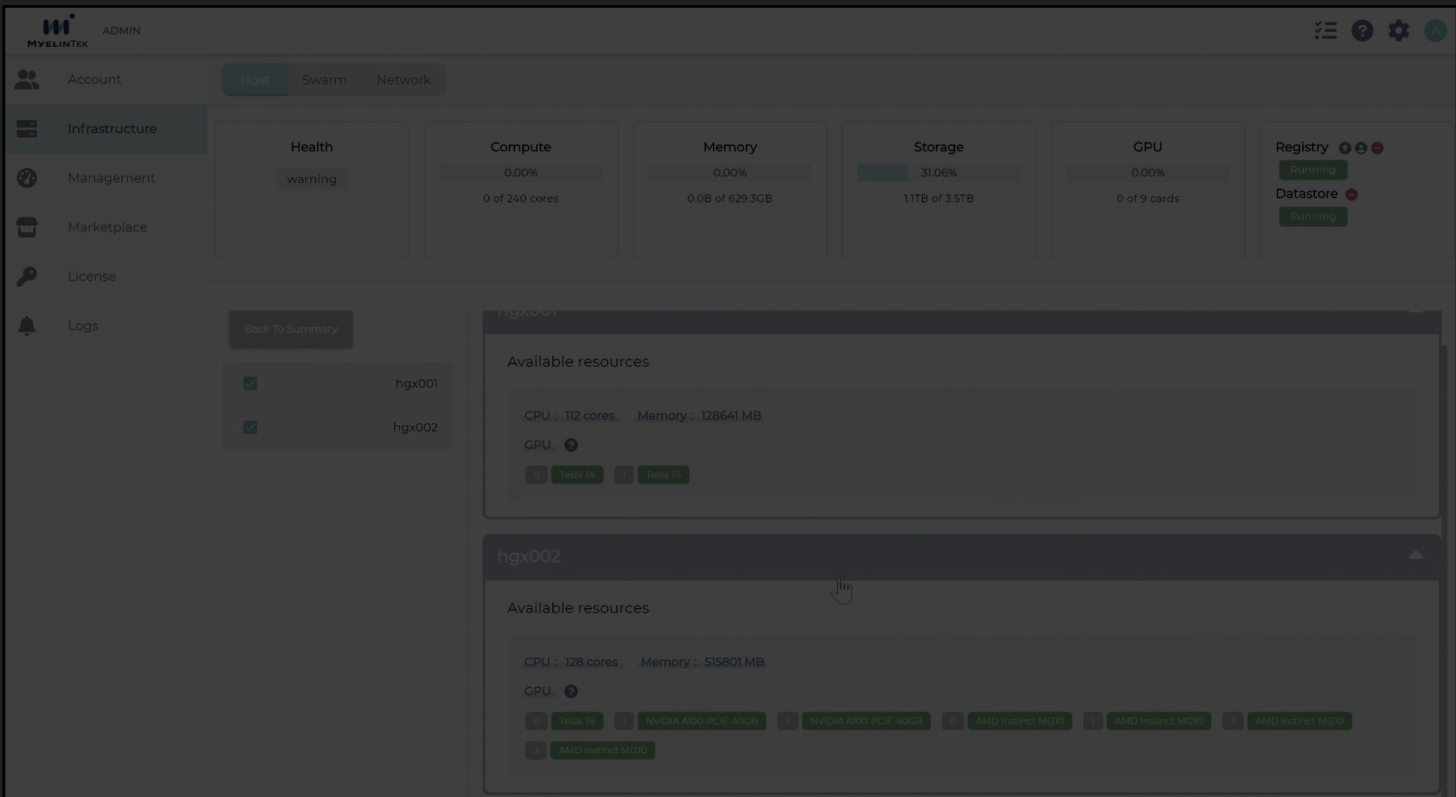
- 訓練參數、過程與產出物版本控制
- 可於平台同時進行多組參數訓練
- 即時顯示資源使用率

模型開發

- 內建常用模型樣板與框架
- 平台介面直接下載NGC及開源資源，如NVAIE、NIM、Huggingface等

管理 – 硬體資源監控與管理

- 平台支援**DELL**全系列伺服器，可納管運算、儲存、管理節點建構資料中心等級的**伺服器叢集**
- 管理者可**自定義資源使用限制**，以供各種模型之硬體配置所需
- 視覺化界面方便辨識既有硬體資源空間與模型專案的使用情況，並可直接關閉容器服務以**回收**相應的硬體資源
- 可以呈現使用者所**佔用的資源**以及其所**佔用的時間**，並產出對應的報表供管理者參考。使用者在的資源使用率可至監控介面個別查詢



Task	Folder	Model	Project	Template	Image	Settings	Backup
Search							
☐	Name	Tag	Summary	CPU Cores	Memory (MB)	CPU (card)	
☐	micro		vCPUs: 2, RAM: 4GB, GPU: 0	2	4096	0	
☐	small		vCPUs: 2, RAM: 4GB, GPU: 1	2	4096	1	
☐	medium		vCPUs: 2, RAM: 8GB, GPU: 1	2	8192	1	
☐	medium-amd		vCPUs: 2, RAM: 8GB, GPU: 1	2	8192	1	
☐	medium-nv		vCPUs: 2, RAM: 8GB, GPU: 1	2	8192	1	
☐	large-amd-1mi210		vCPUs: 8, RAM: 16GB, GPU: 1	8	16384	1	
☐	large-nv-1ai100		vCPUs: 8, RAM: 16GB, GPU: 1	8	16384	1	
☐	large-amd-2mi210		vCPUs: 8, RAM: 32GB, GPU: 2	8	32768	2	
☐	large-nv-2ai100		vCPUs: 8, RAM: 32GB, GPU: 2	8	32768	2	

Username

admin

2025/06/01 12:54

~

2025/07/17 12:54

Total Task Number

7

Total Task Hours

1804.57 hr

Total GPU Hour

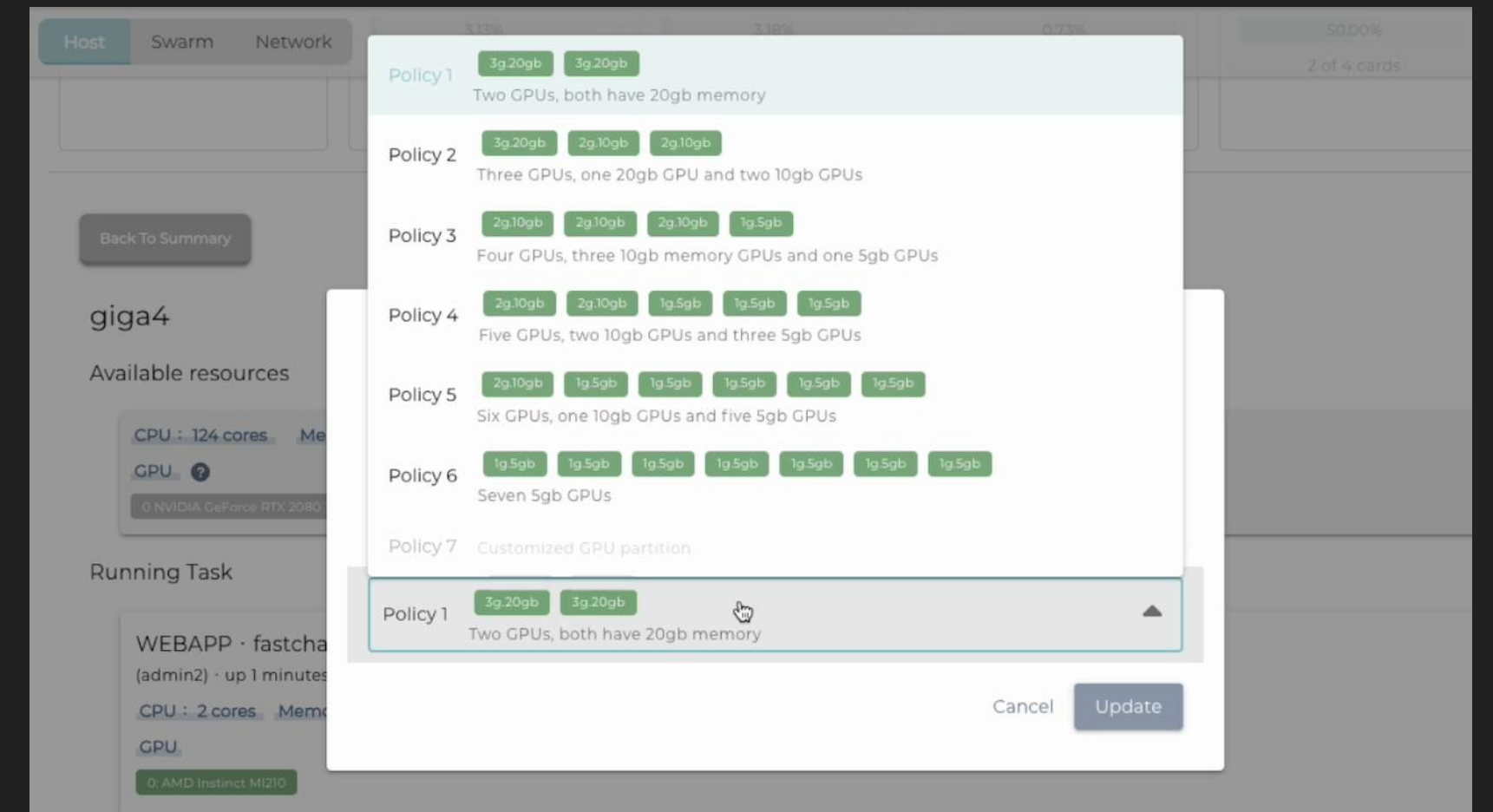
NVIDIA GeForce GTX 1080 Ti - 741.40 hr

Uuid	Username	Total Hours	Flavor	Cpus	Start	End
Ea1da919	Admin	215.57	CPU : 2, RAM : 4096MB, GPU : 0	-	2025/03/14 09:27	2025/06/10 12:28
Ue3d671d	Admin	215.57	CPU : 2, RAM : 4096MB, GPU : 0	-	2025/04/28 12:05	2025/06/10 12:28
Ubb654e8	Admin	741.40	CPU : 4, RAM : 8192MB, GPU : 1	NVIDIA GeForce GTX 1080 Ti X1	2025/06/16 15:30	
Ed04075b	Admin	0.12	CPU : 2, RAM : 4096MB, GPU : 0	-	2025/06/20 10:17	2025/06/20 10:24
E8126814	Admin	76.41	CPU : 2, RAM : 4096MB, GPU : 0	-	2025/06/20 10:25	2025/06/23 14:49
Ub2a399a	Admin	0.04	CPU : 2, RAM : 4096MB, GPU : 0	-	2025/06/23 14:50	2025/06/23 14:52
U49f8100	Admin	555.46	CPU : 2, RAM : 4096MB, GPU : 0	-	2025/06/24 09:27	



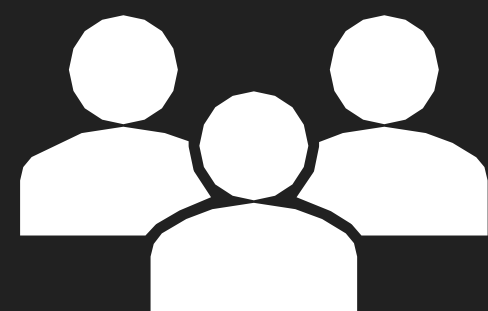
管理 – GPU運算管理

- 將**NVIDIA MIG**功能以下拉式選單操作，快速進行配置
- 平台支援**NVIDIA GPU Fraction**管理功能，提高每單位運算產能
- 平台提供**預留GPU**運算資源功能
- 平台可記錄**GPU運算使用額度**，並產生報表

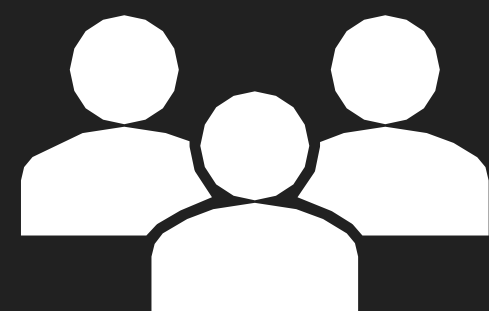


管理 – 專案協同開發

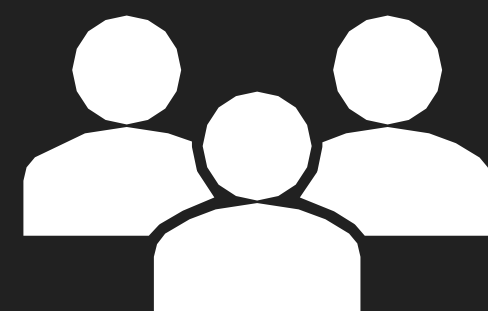
Data Team



AI Team



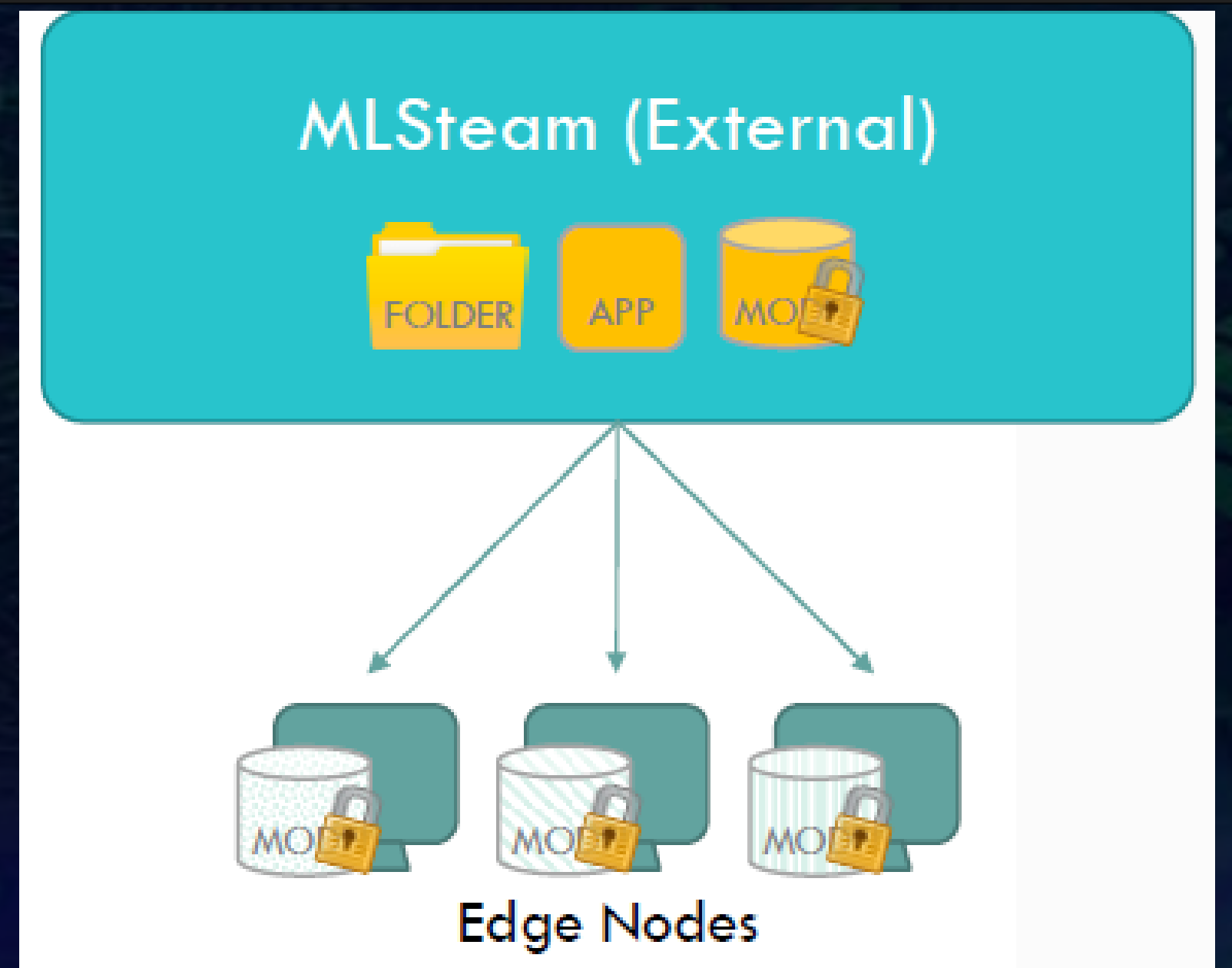
Application Team



 MLSteAM MLOps Platform

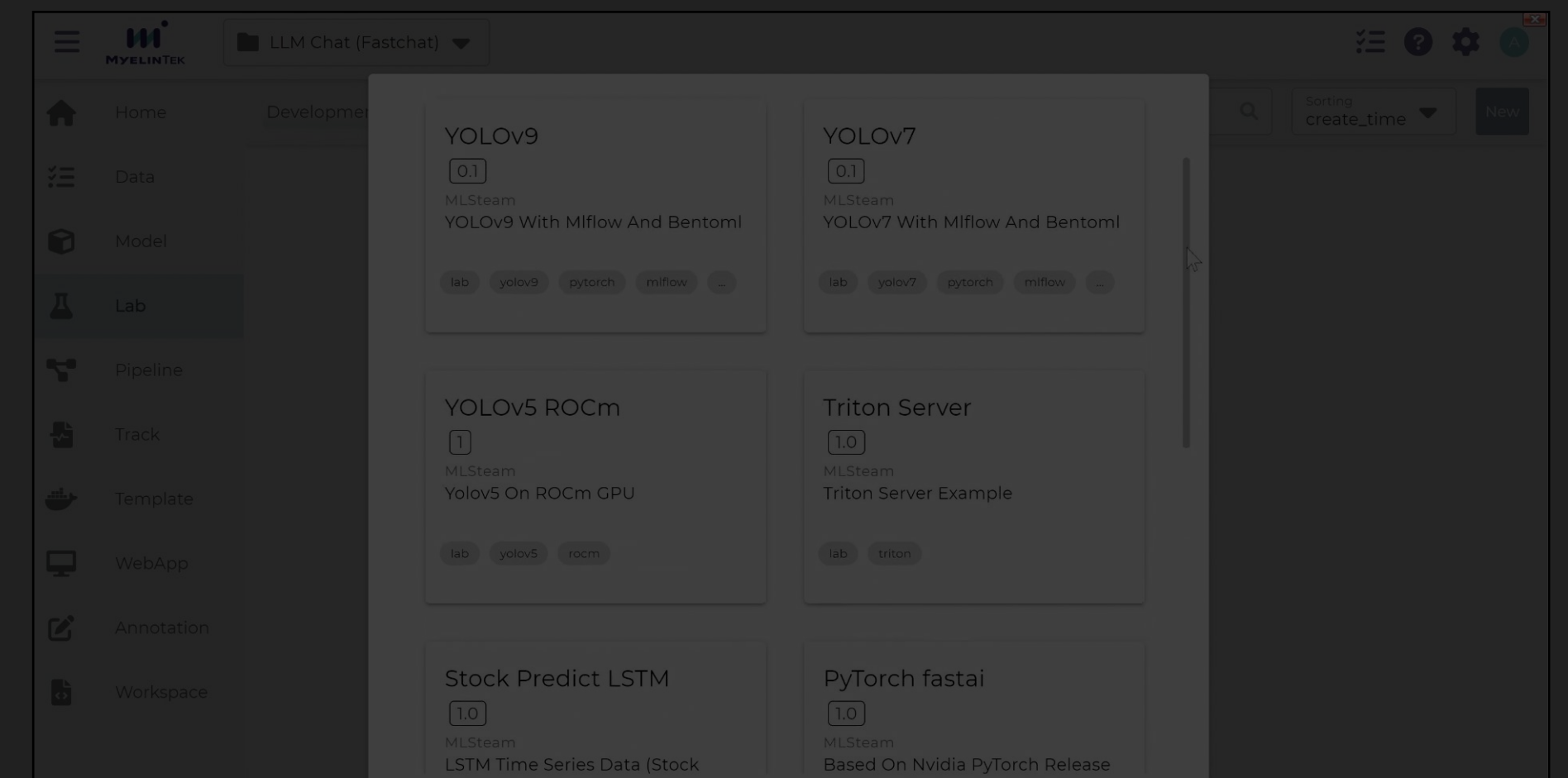
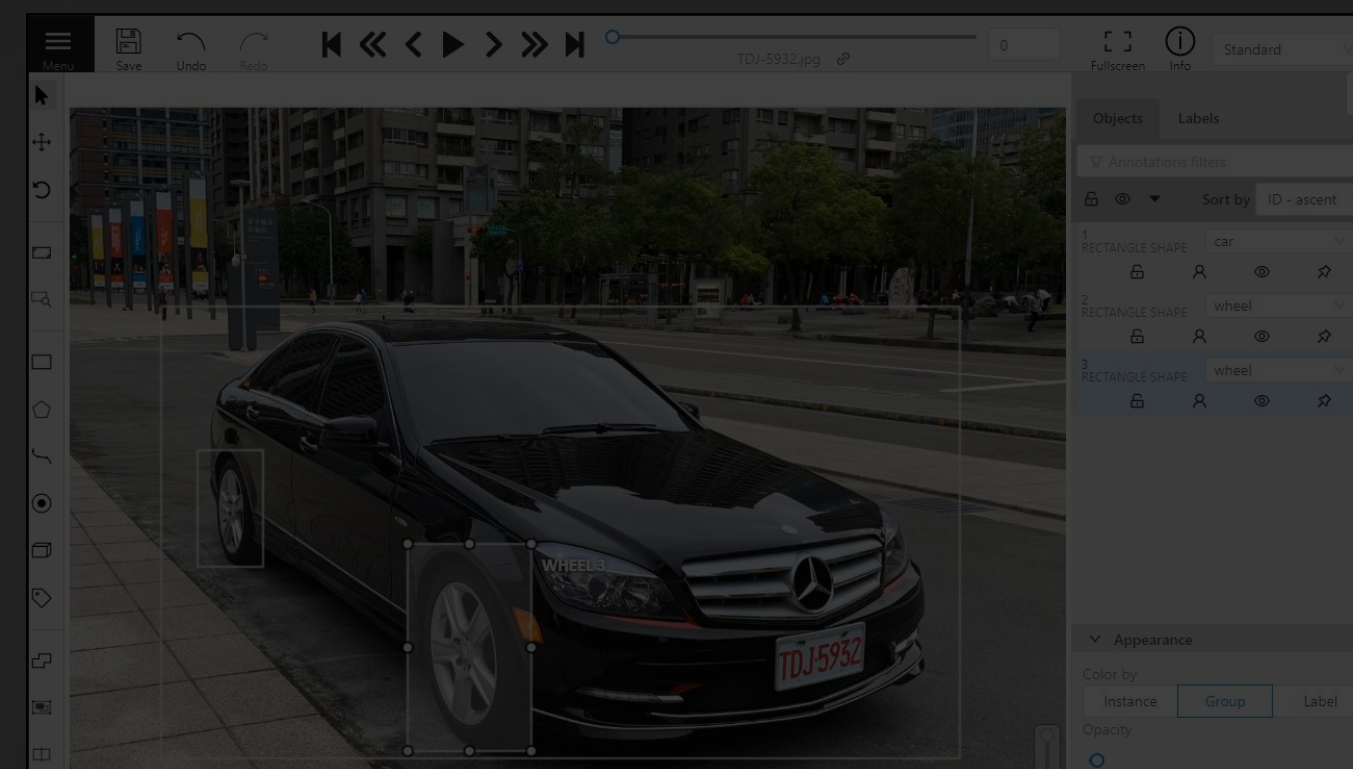
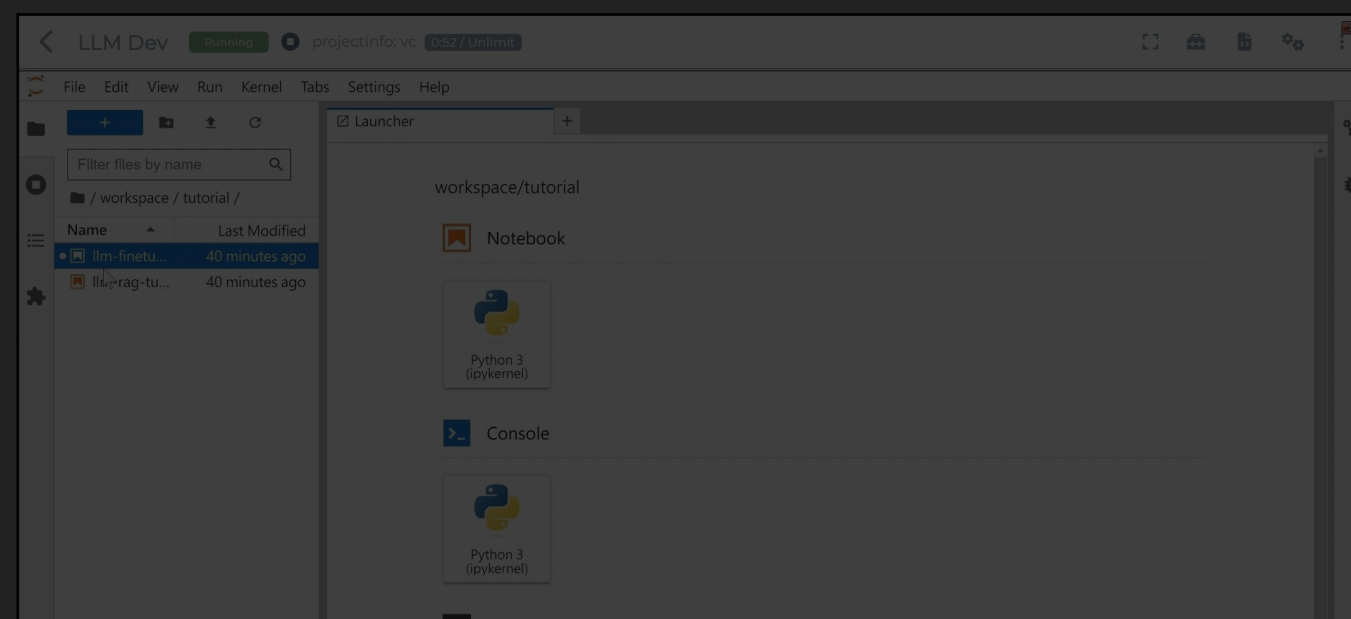
管理 – 平台安全機制

- 模型開發環境與平台介面共用網路埠對外提供服務，無須額外開放網路埠
- 各專案間的獨立作業環境 (Isolation Execution Environment)
- 資料集權限控制管理 (Dataset permission control)
- 模型部署保護機制 (Model deployment protection Mechanism)



整合 – 建構即用的開發環境

- 提供JupyterNotebook與Terminal介面的模型開發環境
- 平台支援PyTorch、TensorFlow等深度學習框架
- 平台可支援NVIDIA CUDA、AMD ROCm、Intel OneAPI等異質運算
- 平台內建CVAT、LabelStudio等標註工具
- 預建Llama、YOLO等常用之模型樣板，其版本隨平台進行更新
- 容器啟動可透過GUI設定SSH port，不用編輯yaml



整合 – 支援NVAIE與NIM

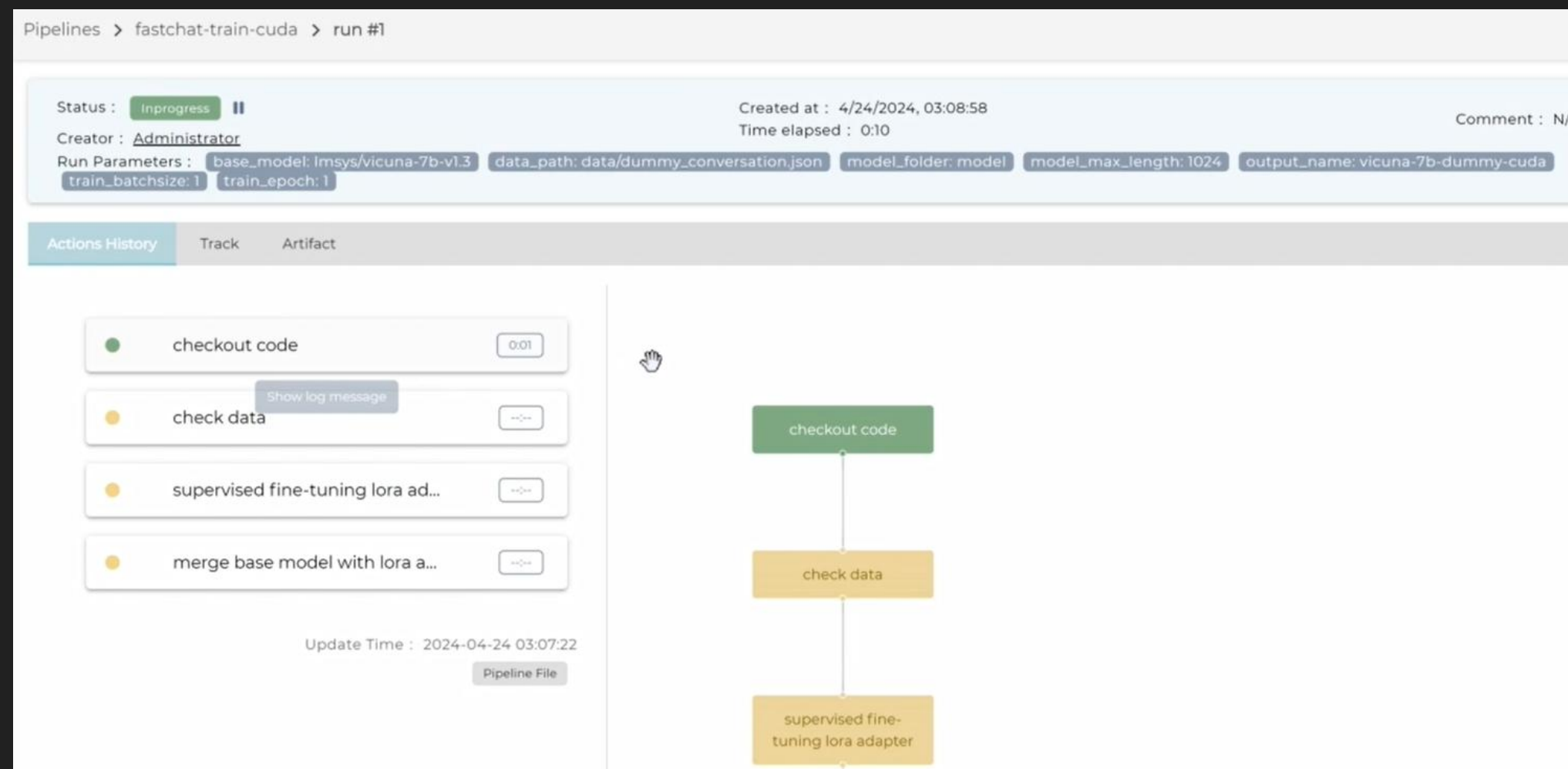
- 平台支援NVAIE模型開發工具及框架，如：NeMo、RIVA、MONAI...
- 透過平台介面，易於使用NVIDIA NIM(微服務)

The screenshot displays the MYELINTEK platform interface, which is a Jupyter Notebook environment. The notebook is titled "End-to-End Experimentation with NVIDIA NeMo" and shows a "Settings" panel with various configurations. The code in the notebook includes imports for NeMo and PyTorch, and a function to process audio files. The interface also shows a "Lab Resource" panel with metrics for Compute (41%), Memory (29%), Storage (0%), and GPU0 (0%).

On the right side of the image, the NVIDIA NGC Catalog is shown. It features a search bar and a list of containers, including NVIDIA NIM and NVIDIA Enterprise Platforms. The catalog also displays a "Containers" section with a description: "The NGC catalog hosts containers for AI/ML, metaverse, and HPC applications and are performance-optimized, tested, and ready to deploy on GPU-powered on-prem, cloud, and edge systems."

整合 – Pipeline/MLFlow功能

- **CI / CD整合**，助於持續性的模型訓練與優化
- 直觀且視覺化的進度顯示與日誌紀錄
- 依照資源需求進行調度，將硬體資源利用最大化
- 使用YAML檔案輕鬆**定義執行相依性 (DAG)**
- 透過**Pipeline**定義執行內容，並安排排程讓服務在指定期間啟動，執行完畢後**自動關閉將資源釋放**，以進行批次推論服務

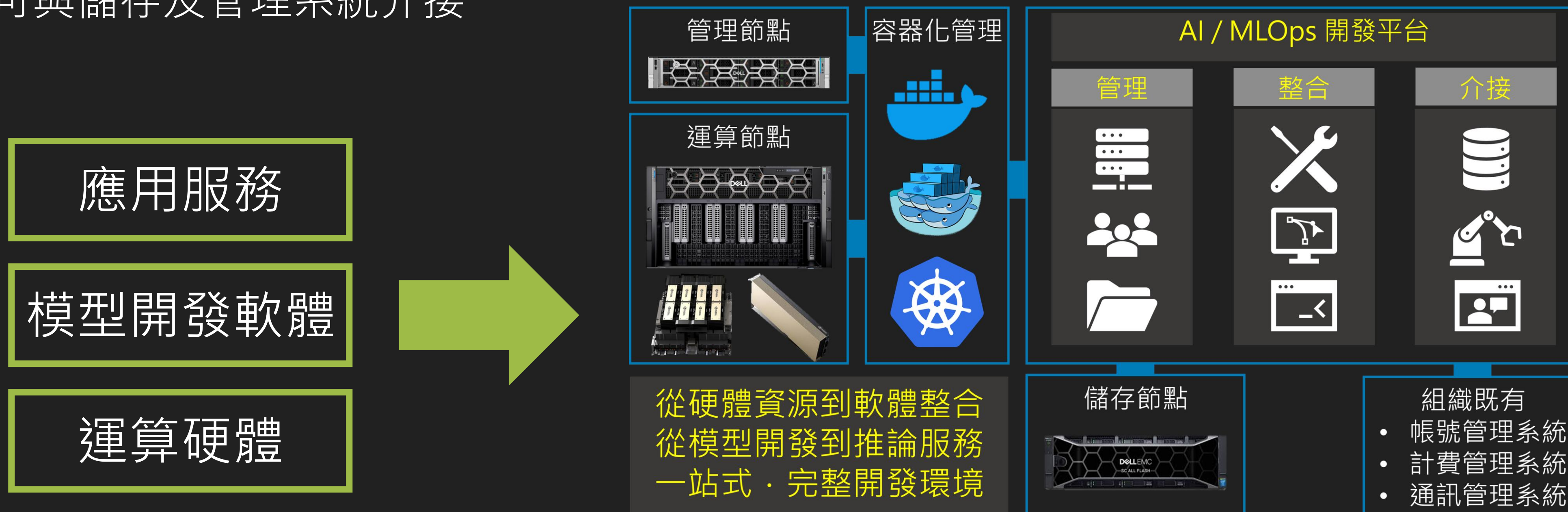


介接

- 平台支援LDAP等管理系統
- 平台可與儲存系統進行介接
- 管理者可透過平台的quota管理功能與既有計費系統進行介接
- 平台可介接端點系統 (edge伺服器、AIoT)，進行模型部署、版本管理、資料收集等
- 可於平台啟動WebApp進行模型推論服務

總結 – 邁爾凌AI Platform

- 工研院spin-off的擁有多多年AI開發團隊
- 安裝於地端的一站式AI開發平台
- 成功案例：電信商、國防、智慧交通、輔助醫療等...
- 平台所提供三大項目：
 - **管理** → 使用者、硬體、資料集、模型版本
 - **整合** → 開發所需工具，如標註工具、開源模型、框架並支援NVIDIA NVAIE、NIM
 - **介接** → 可與儲存及管理系統介接



Action Now

Brad Huang 黃丞毅

brad@myelintek.com

+886-933-306-785

