



HPE GREENLAKE : 掌握數位轉型新趨勢

Janet Lin / GreenLake Cloud Business Development
Nov, 2021

下一波
數位轉型浪潮
將由無所不在的應用程式和資料來推動

數位轉型僅涵蓋了一小部分的應用程式和資料

可輕鬆
移至雲端的應用程式，現已完成移轉

現代化雲端體驗
速度與敏捷性

內部部署實情
資料引力、安全性、
效能、合規性

70%
的應用程式仍位於公有雲之外¹

1 IDC Cloud Pulse 2020 年第 1 季

僅限HPE內部使用 | 3

現在，數位轉型可全面提升速度

將雲端體驗應用於
所有應用程式和資料 —
隨時隨地

客戶面臨的挑戰

業務工作負載需要仰賴雲端經驗所能發揮的**敏捷度**，
才能戰勝競爭對手

組織需要刪減**成本**，
節約資本，
並且依據業務需求調整成本

數位業務需要搶先需求
一步進行快速**擴充**

多雲端的混合 IT 作業相當
複雜，會提高成本、
風險並拖慢獲利速度

HPE GREENLAKE

為您量身打造的雲端服務



自助服務

依使用量付費¹

擴大/縮減
規模

代您管理

¹ 可能必須有預留空間

僅限HPE內部使用

藉助 HPE GREENLAKE 加速實現價值

時間

75%

縮短部署數位專案的時間¹

風險

85%

意外停機時間減少²

成本

30%-40%

資本支出節省，因為不再需要過度佈建¹

控制與洞察

40%

減少 IT 的支援負荷，進而提高 IT 團隊的生產力¹

1 委託 Forrester Consulting 進行的研究，《HPE GreenLake 的整體經濟影響》(The Total Economic Impact™ of HPE GreenLake)，2020 年 5 月

2 由 HPE 委託製作的 IDC 白皮書，《HPE GreenLake 管理服務的商業價值》，2020 年 1 月

HPE GREENLAKE 雲端服務



快速上手，輕鬆入門

符合業務規模所需—入門門檻不高，能隨著您的發展步調獲得成長
隨時能應用於要求最嚴苛的工作負載

價格透明



[HPE.com/greenlake](https://www.hpe.com/greenlake)

最佳化及標準化的組態

超低成本、
最佳效能、
完美平衡
(小型/中型/
大型)



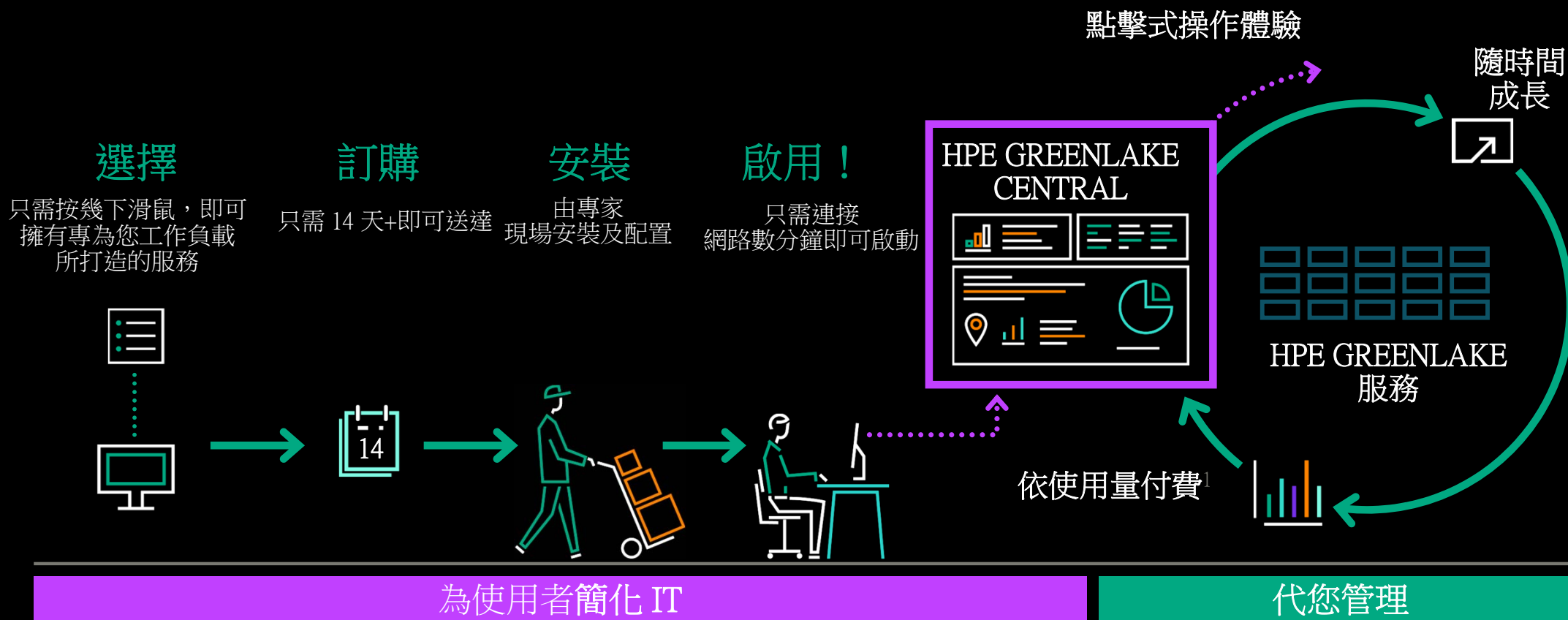
硬體/軟體預先設定且最符合工
作負載需求的，
14 天+即可送貨到府

自助式
雲端體驗



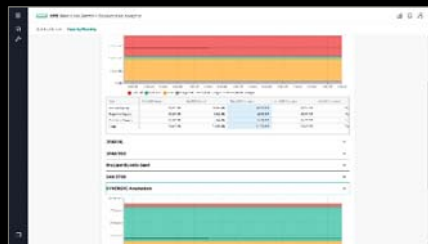
點擊式平台方便您
學習、定價及操作

HPE GREENLAKE 可加快您享受雲端優勢的速度



¹ 可能必須有預留空間

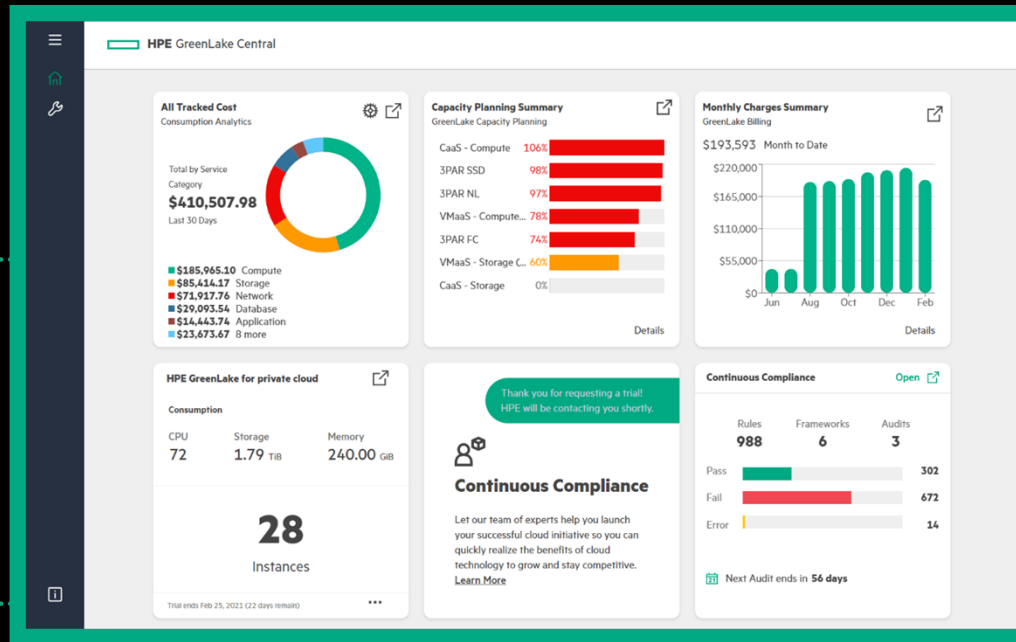
透過 HPE GREENLAKE CENTRAL 重新定義體驗



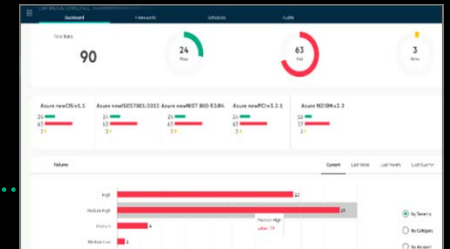
透過內建緩衝區提前預備需要的容量



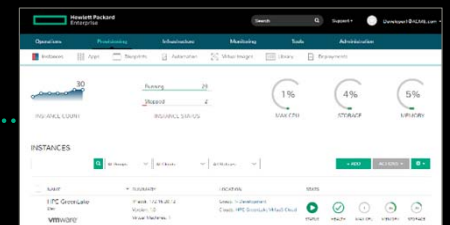
進行成本控制以最佳化雲端支出



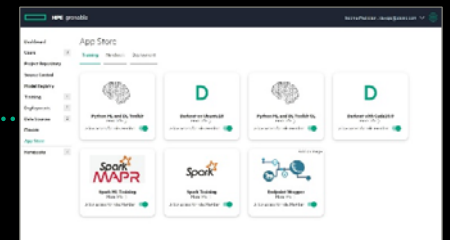
HPE GreenLake Central
直覺、自助和點擊式
入口網站及管理主控台



針對超過 1,500 個控制項
提供持續的合規性



使開發人員能快速輕鬆地
運用資源



透過各種工具選項，更快地為資
料科學團隊提供精闢見解

僅限HPE內部使用 | 11

可協助您加快數位轉型的專業知識

專業知識

邊緣到雲端數位轉型的專業知識，能幫助您想像、設計以及實作合適的人員、流程和技術方法，將雲端經驗擴展至整個 IT 環境

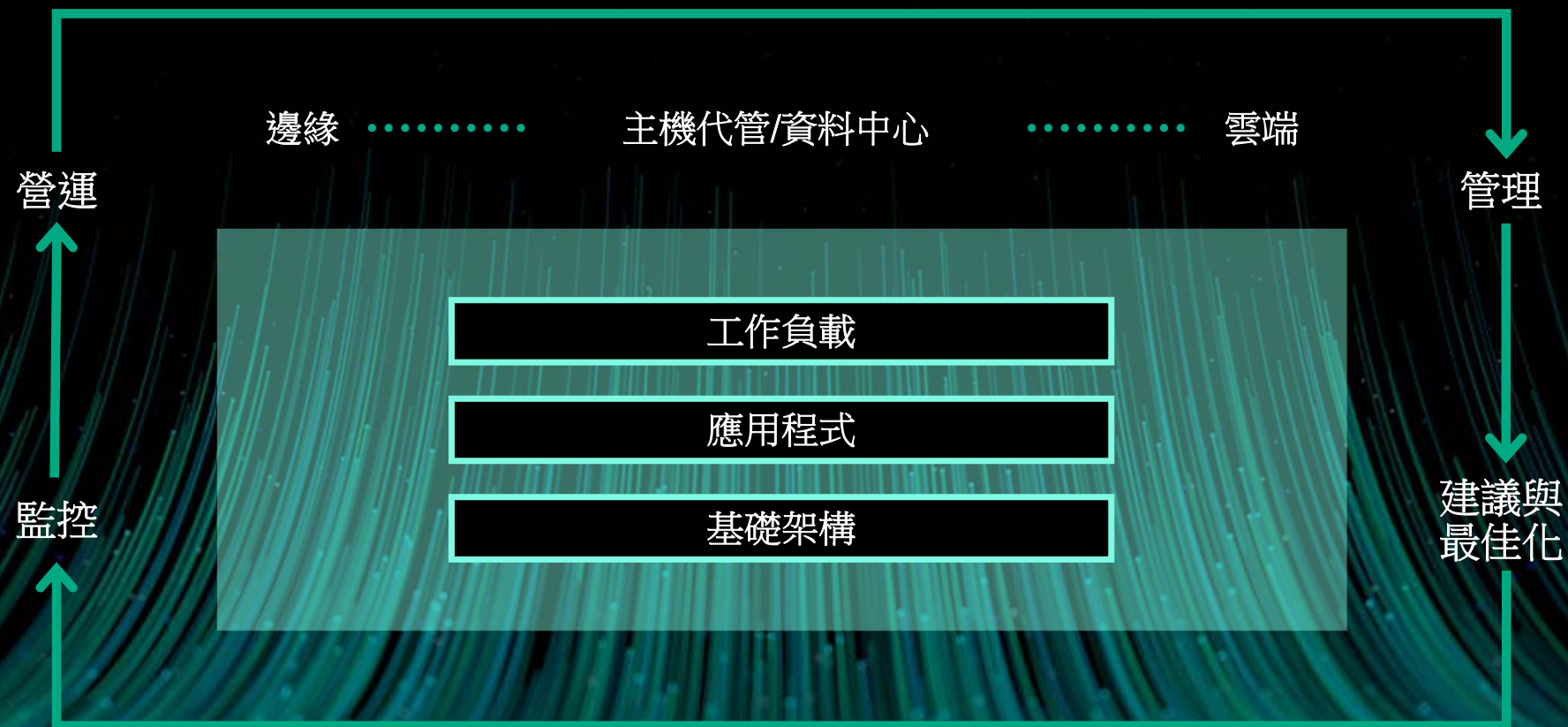
方法

混合雲端採用架構以數千次雲端合作中獲得的經驗為基礎，幾經實證且備受肯定，能協助您評估自己在混合雲端轉型過程中的定位，並且排定適當措施的優先順序

工具

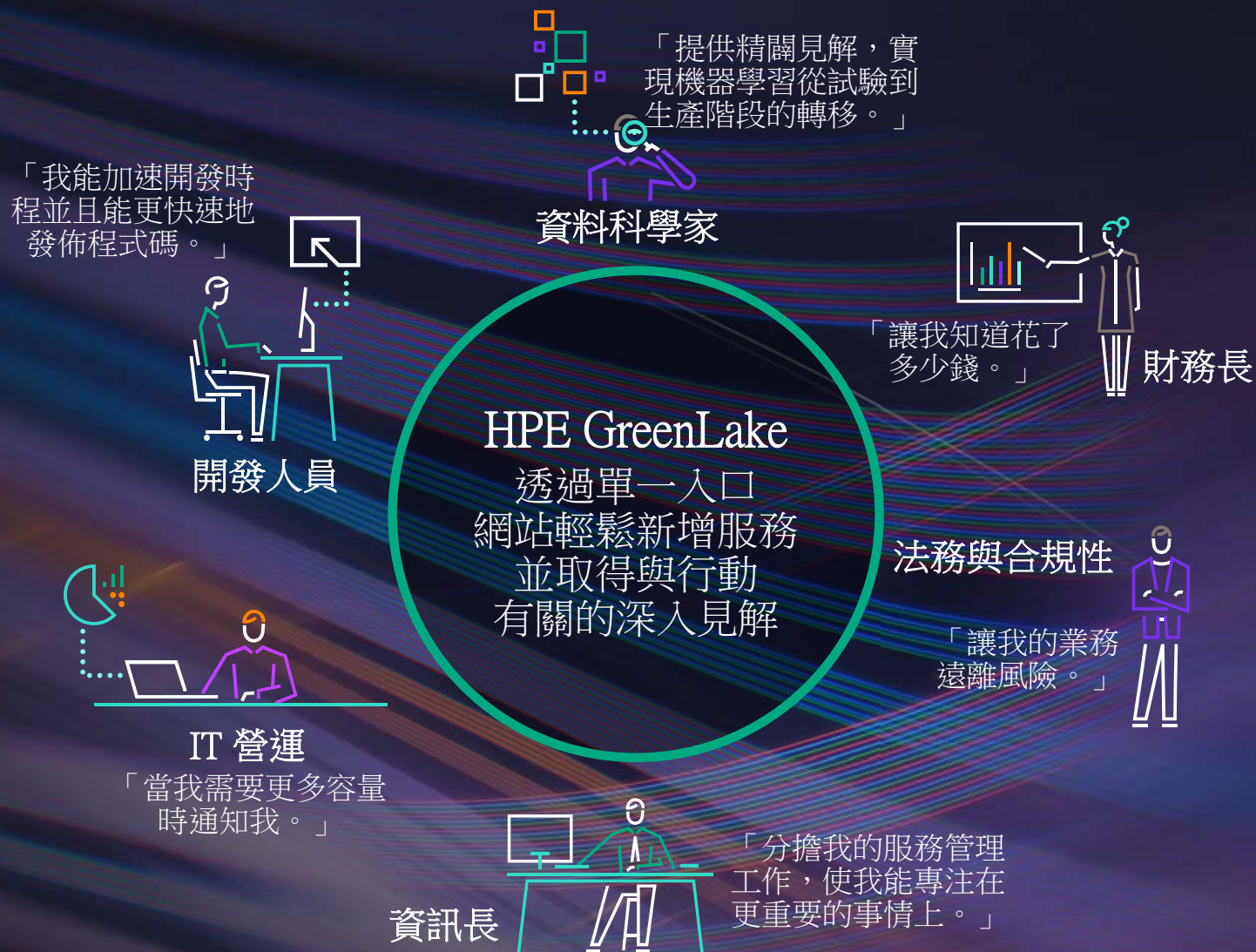
自動服務全面整合於 HPE GreenLake，能運用整個 IT 資產所產生的資料，讓您清楚掌握成本、合規性，以及正確的應用程式與資料混合工作負載安排

HPE GREENLAKE 管理服務



實現服務 交付轉型

加快業務範圍內的服務交付
速度



運用 HPE GREENLAKE 推動永續發展

適當調整 IT 資源
並提高使用率

減少消耗能源
和成本

採用最新、
最有效率的技術

合理淘汰老舊資產，同時為日
後創新供應所需資金

30%
整體持有成本節
省比例¹

30%
能源成本節省比例²

- 1 Forrester, HPE GreenLake 對總體經濟造成的影響 (Total Economic Impact of HPE GreenLake), 2020
- 2 HPE 內部計算結果

HPE GREENLAKE 的 領先優勢

獨特

實惠

部署依使用量付費¹，
以實際使用量計算結果
為準

有效

容量規劃；
節省成本並避免過度佈建

專業

協助客戶進行設計、交
付、營運

專業知識

45 億美
元

截至目前為止，
預定的 TCV

約 1000 家

客戶，涵蓋全球 50 多
個國家/地區

專有 IP

透過收購和合理化投資實現

800 多個

合作夥伴隨時
服務和銷售

12 年以上

IT 即服務的
部署經驗

95%

留客率

← 具有深度和廣度的
產品、合作夥伴、技術 →

¹ 可能必須有預留空間

後續步驟

- 評估並確定您的最佳雲端組合
- 請觀看 HPE GreenLake 示範
- 運用 [Forrester Research ROI 估算工具](#)，瞭解使用 HPE GreenLake 的投資報酬率
- 立即透過 HPE GreenLake 試用及購買雲端服務

如需更多資訊，請造訪 hpe.com/GreenLake

謝謝

Janet Lin; janet.lin@hpe.com





Hewlett Packard
Enterprise

SOLVING AI PRODUCTION DEPLOYMENT CHALLENGES

沈仲杰

Nov 2021

AI & ADVANCED ANALYTICS ARE TOP PRIORITIES FOR ENTERPRISE

- IDC and Gartner show that 2 of the top 3 priorities of CIOs are focused on AI and advanced analytics.
- 15 to 20% of all new infrastructure deployed will be used to support these workloads by 2021.
- AI/ML/DL is projected to continue to be one of the hottest, highest growth areas in enterprise adoption.

¹ IDC, Goldman Sachs, HPE Corporate Strategy, 2018

² Gartner - "2019 CIO Survey: CIOs Have Awakened to the Importance of AI"

CIO Priorities

2/3

of top priorities^{1,2}

Infrastructure

20%

usage in 2021¹

Adoption growth

2.7X

over 4 years²



DEPLOYING AI HAVE OPERATIONAL CHALLENGES



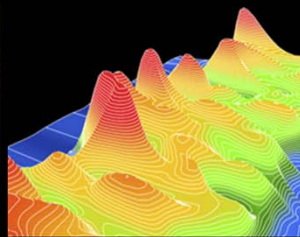
Financial services

Fraud detection,
ID verification



Government

Cyber-security, smart cities,
defense, and utilities



Energy

Seismic and reservoir
modeling



Retail

Video surveillance,
shopping patterns



Health

Personalized medicine,
image analytics



Consumer tech

Chatbots



Service providers

Media delivery



Manufacturing

Predictive and prescriptive
maintenance

80–85%

of enterprises are running into the 'last mile' problem with ML model deployment and management.¹

73%

of enterprises fail to deploy their AI projects into production. The reasons these projects fail have little to do with having the best data science team and more to do with standard business and IT integration challenges.²

¹Sumit Pal, Sr Director Analyst, Gartner, "Don't Stumble at the Last Mile: Leveraging MLOps and DataOps to Operationalize ML and AI"

²DC, Market Analysis Perspective, Worldwide Artificial Intelligence Software, September 2020

TOP REASONS WHY ENTERPRISE AI PROJECTS FAIL TO MAKE IT INTO PRODUCTION

Solving the wrong problem and not scoping the AI project correctly

- Applying AI towards incorrect problems
- Setting in proper expectations
- Using AI to solve a problem with different analytical techniques
- Ignoring requirements in order of magnitude of resources across multiple business groups
- Not utilizing resources beyond data science teams when creating and refining models

Not securing the necessary funding

- Overbooking that AI projects are expensive and resource intensive
- Failing to identify the right problem and scope to properly fund
- Not acquiring the right resources to fully scale an AI project

ML Ops does not equal DevOps

- Not integrating ML Ops into how enterprise IT environments operate their DevOps
- Neglecting how ML models are trained and refined for specific AI use cases
- Forgoing AI hardware that is compliant, secure, and reliable to meet enterprise IT requirements

Data pipelines not ready for real-time data

- Neglecting how real-time data is used and passed off into training, inference, testing, and deployment
- Not selecting technologies that will be used to funnel the data to the right destinations
- Overbooking how real-time data will be cleaned and prepped before being used in an AI use case
- Ignoring how data is collected and sent to the data center (batch versus stream)
- Not considering real-time data streaming from multiple sources and how it gets plugged into models

IT infrastructure not ready for AI

- Disregarding that enterprise data centers must have AI-ready hardware and software
- Not realizing certain processors and tuned software is needed to run AI algorithms efficiently
- Ignoring how to operate high-powered systems with GPUs to accelerate performance
- Not setting up for storage and memory systems that can handle fast data transfers
- Not utilizing the right infrastructure for generating and sending massive amounts of data from the edge

HPE BEST-IN-CLASS TECHNOLOGY WHERE AND WHEN YOU NEED IT



Location

Edge

Core

Cloud



Workloads

Analytics

Training

Inference



Infrastructure

Compute

Storage

Networking

Software



Sourcing model

Purchasing

Financing

As-a-service



Accelerated compute

FPGA

GPUs

Memory-driven compute



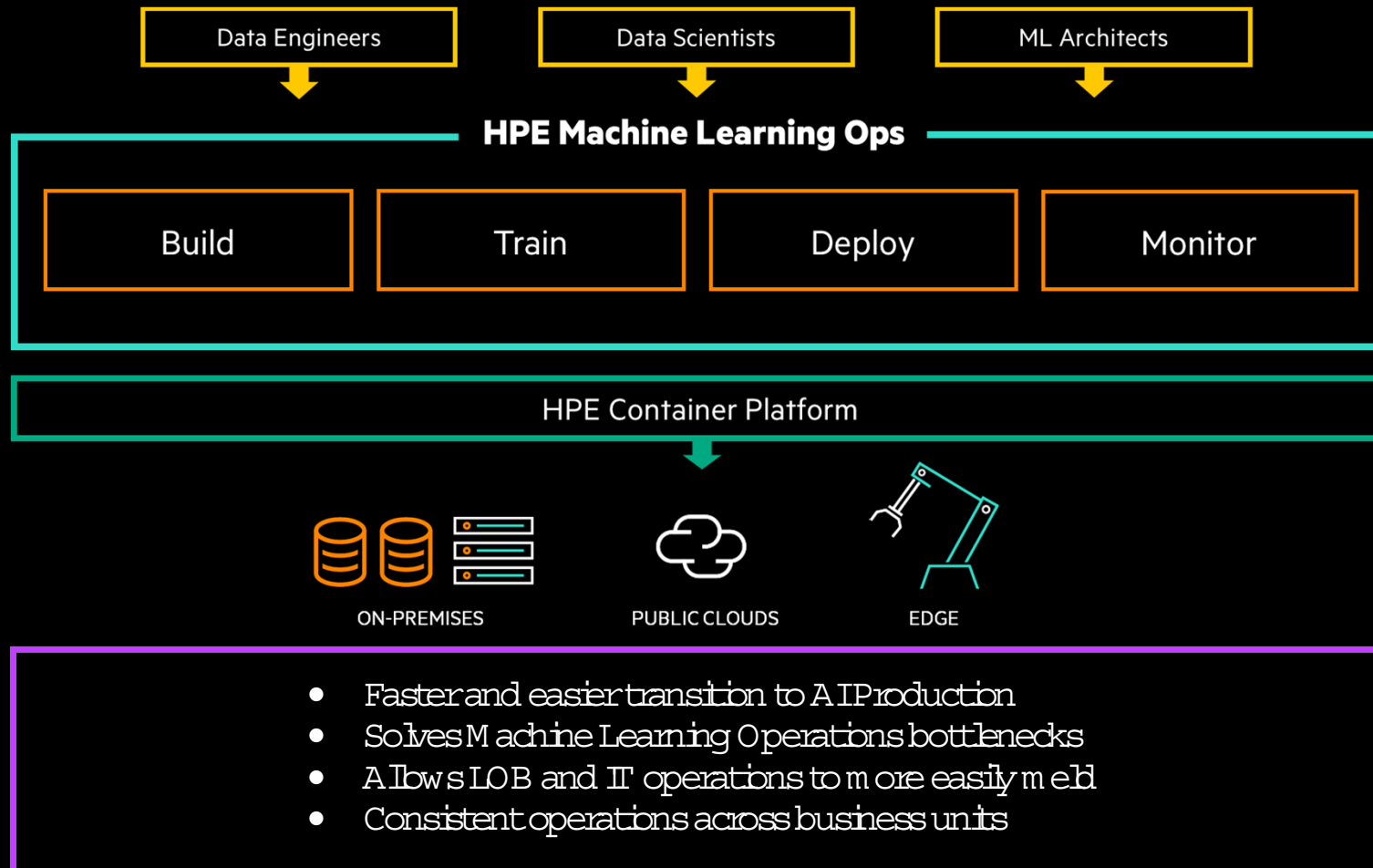
Deployment stage

POC

Scaled POC

Production at scale

ACCELERATE TIME-TO-VALUE FOR AI WITH HPE ML OPS



FULL COMPLEMENT OF ENTERPRISE AI-READY ARCHITECTURE CHOICES

AI Ready Architecture

HPE Emerald

HPE Container Platform and MLOps

NVIDIA GPU Cloud

HPE supported NGC AI/ML libs

Do-It-Yourself

Bare-Metal, VM, Containers

Cray Unika

Supercomputing Software Platform

Choice of AI/ML Libs— TensorFlow, Caffe, PyTorch, MXNet, TF Serving, Clipper, Flask, ONNX

Accelerators— CPU (Intel/AMD), GPU (NVIDIA/Intel/AMD), FPGA (Intel/Xilinx), ASIC (Intel)

Fabric— Ethernet (Aruba, Slingshot), InfiniBand (Mellanox)

Storage— MapR (ApoIb/ProLiant/EdgeLine), ClusterStor (E1000), Weka (ApoIb/ProLiant)

Compute— ApoIb 6500/2000, Cray Supercomputers, Superdome, ProLiant DL380/385, EdgeLine EL1000/4000/8000