

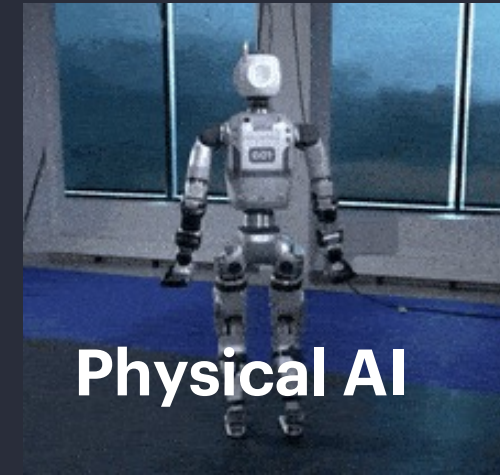
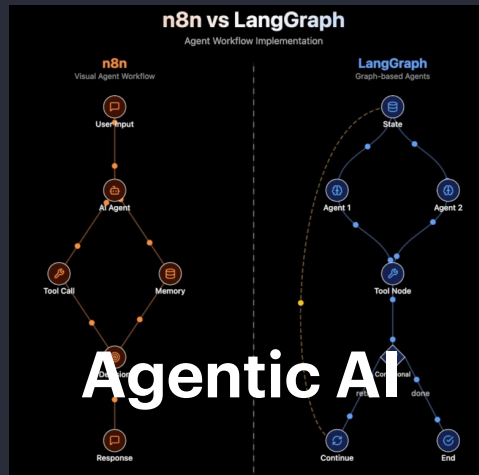
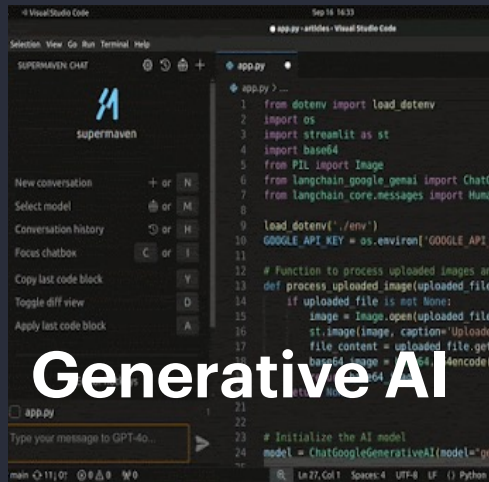


Zielgerichtet KI im Mittelstand nutzen

Jeremy Massow

April 2026

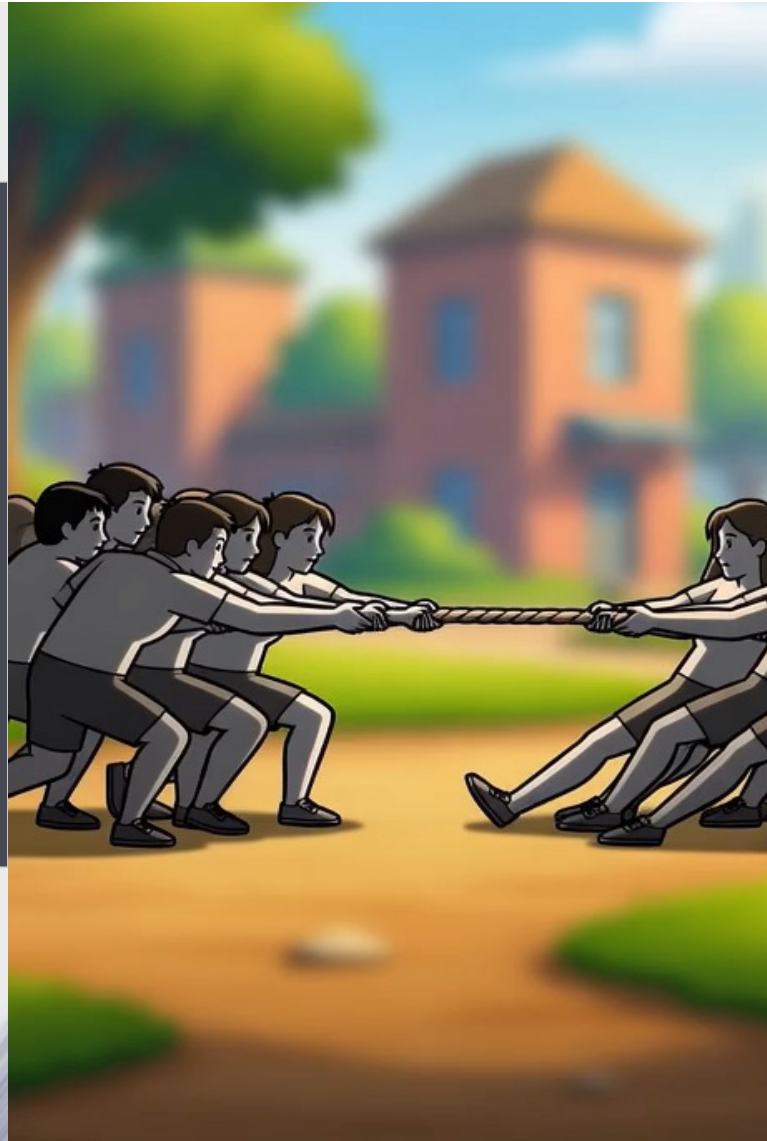
Die drei Bereiche für AI-GETRIEBENE TRANSFORMATION



FOMO

[Fear of Missing Out]

- Wettbewerbsdruck
- Hype-getriebene Dynamik
- Ecosystem lock-in
- Kurzfristige Erfolge



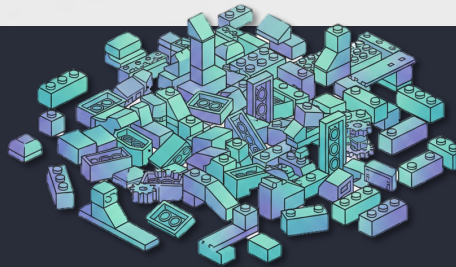
FOMU

[Fear of Messing Up]

- Integrationsrisiken
- Unsicherheit zum ROI
- Regulatorische und ethische Bedenken
- Überlastung

Al am Beispiel eines

LEGO Raumschiffs



DIY

Benötigt viele Bausteine, umfassendes Wissen über Lego & Star-Wars, Zeit, Trial & Error, maximale Komplexität und Risiko

oder



Buy the kit

Bausteine und Pläne sind beinhaltet, jedoch 1.949 Einzelteile und 239 Seiten Anleitung. Signifikante Zeit und Aufwände notwendig

oder



Buy the finished product

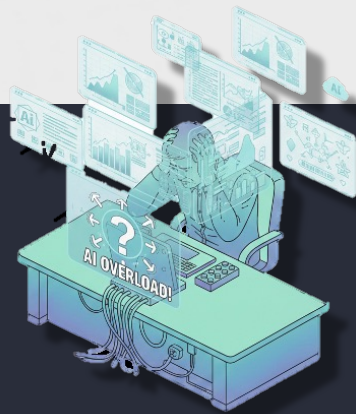
Alles ist spielfertig vorbereitet und einsatzbereit.



AI Implementierung

Der Grad an Komplexität

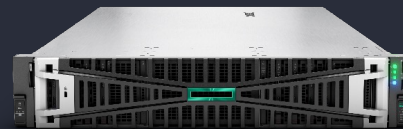
Laut MIT-Studie steigt die Erfolgsquote auf bis zu 67% wenn Unternehmen auf externe / vorgefertigte Lösungen statt Eigenentwicklungen setzen



DIY

Erfordert umfassende Infrastrukturkenntnisse, umfassende Softwarekenntnisse, Zeit, maximale Komplexität und Risiko

oder



Buy the kit

Vorkonfigurierte Infrastruktur mit teilweise installierter Software. Noch umfangreiche Arbeiten erforderlich, um Lösung einsatzbereit zu machen.

oder

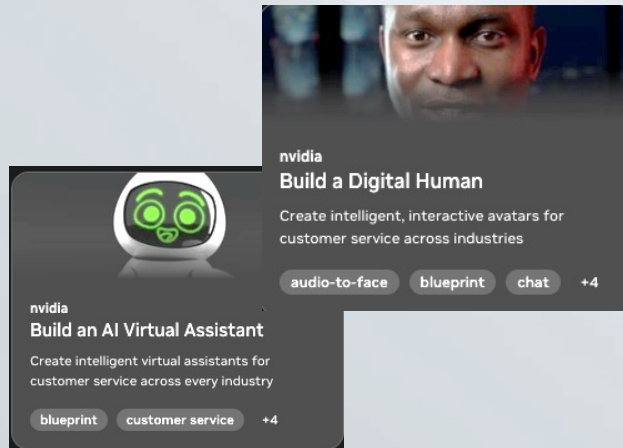


Buy the finished product

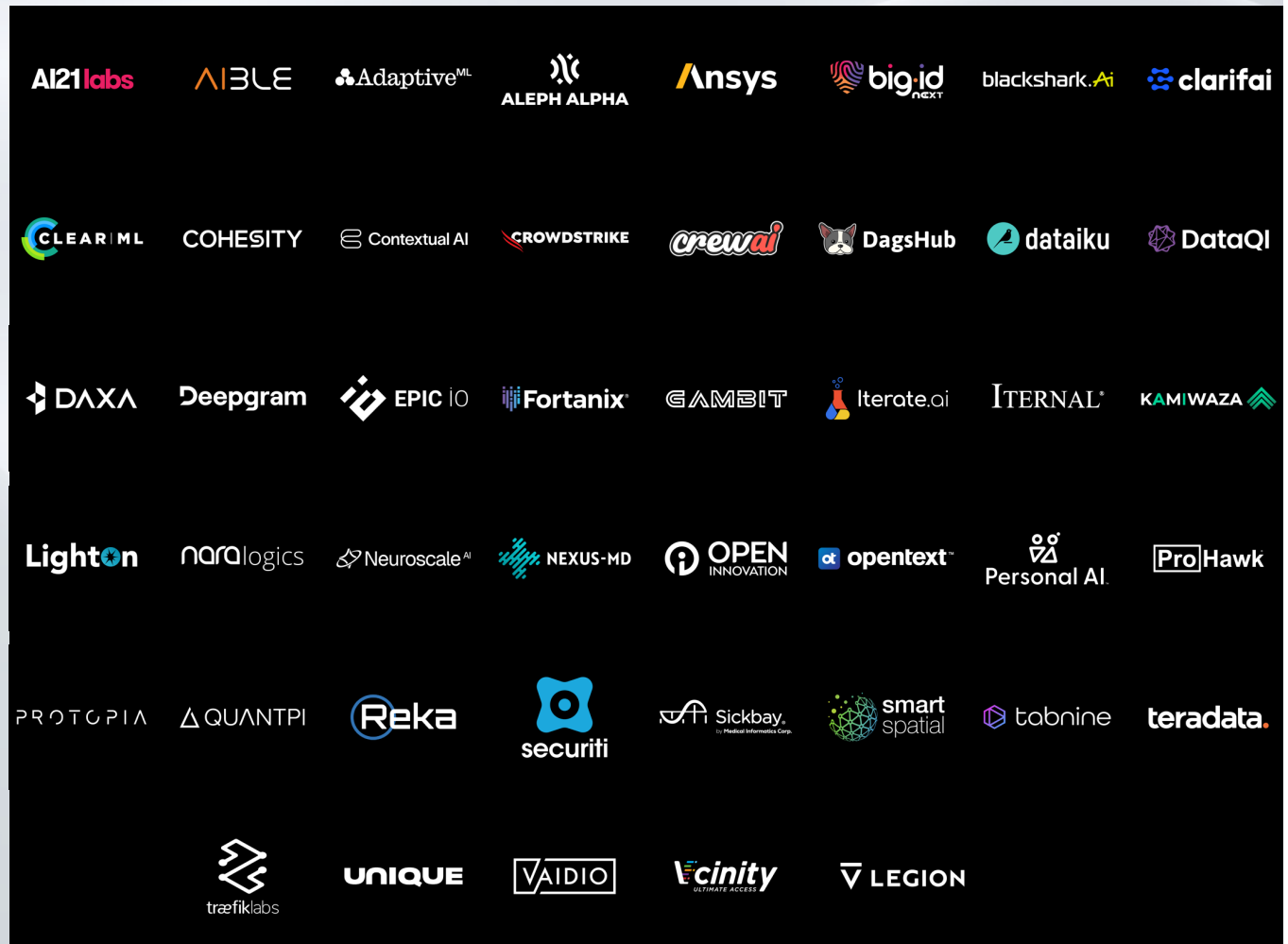
Vollständig integrierte Lösung, mit minimalem Einrichtungsaufwand.. Ideal, wenn Sie Ergebnisse statt Komplexität wünschen.



Unleash AI ISV Partner Program



NVIDIA Blueprints



In 3 Leveln zur AI-Mastery

3.

AI Mastery durch die Nutzung von AI Plattformen:

Zentralisierung der AI Kompetenzen durch den Aufbau einer holistischen AI Plattform wie **Private Cloud AI** und dem **AI Factory Portfolio**.

2.

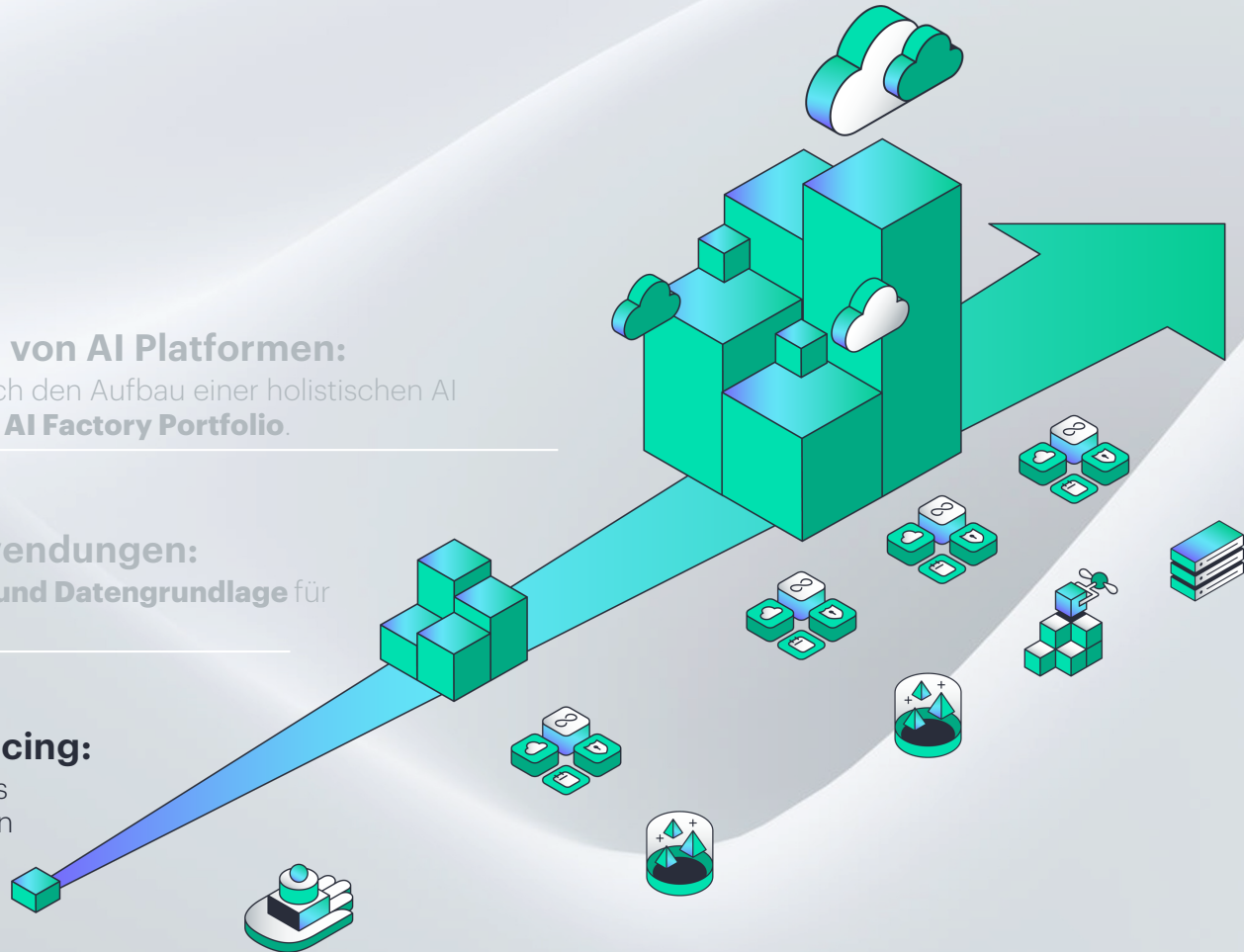
Networking & Data für AI Anwendungen:

Schaffung einer **AI-fähigen Netzwerk und Datengrundlage** für zukünftige AI Anwendungen.

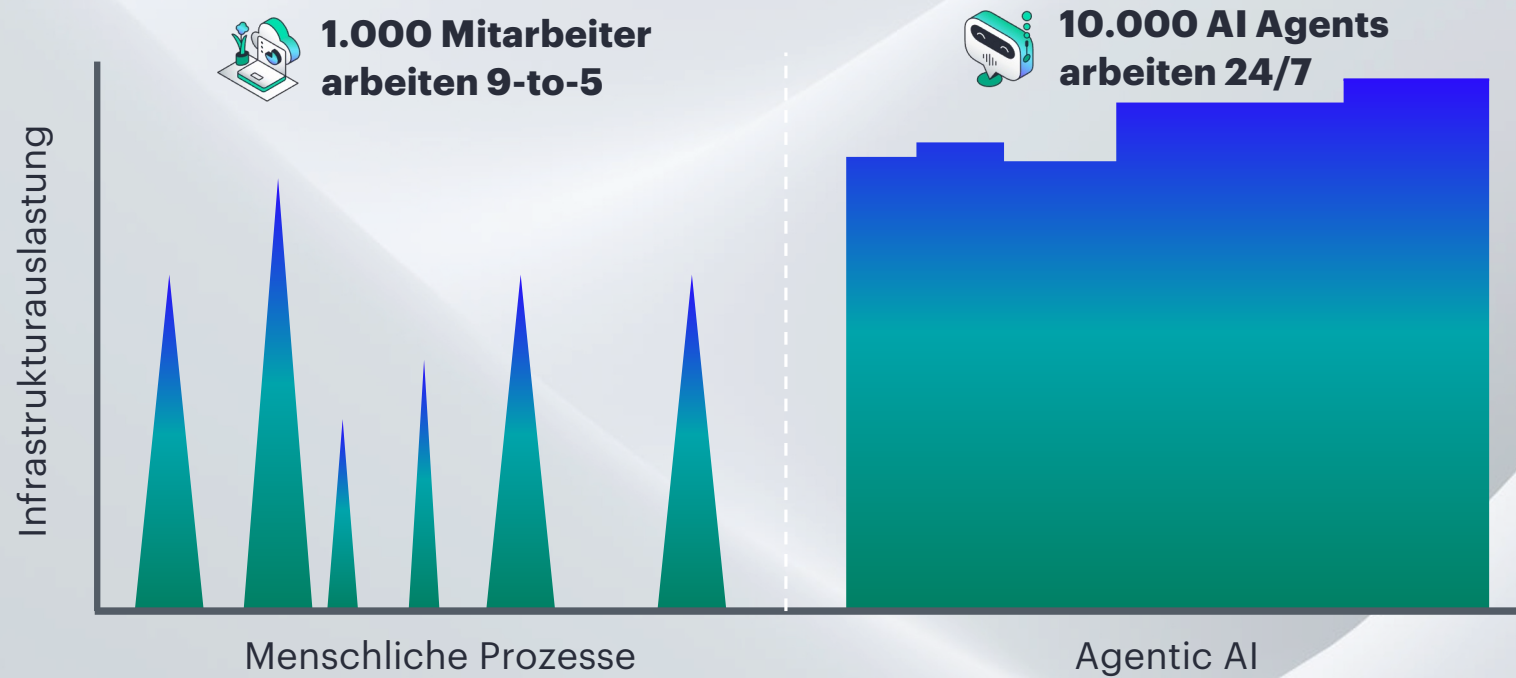
1.

Rechenzentrum für AI Inferencing:

Nutzung des KI-fähigen Server Portfolios **inkl. 3rd party SW** und Anforderungen an Rechenzentren.



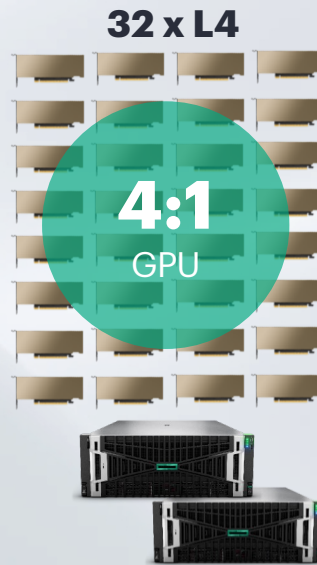
Inferencing außerhalb der “Business Hours”



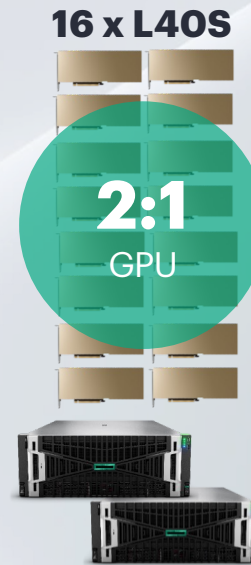
GPU Performance Sprünge sind massiv

Bis zu:

5X mehr Performance
48% niedrigere Kosten
25% Energieeinsparung

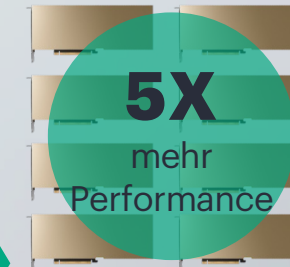


4 X DL380a Gen12



2 X DL380a Gen12

8 x RTX PRO 6000



1 X DL380a Gen12

Kostensparnis

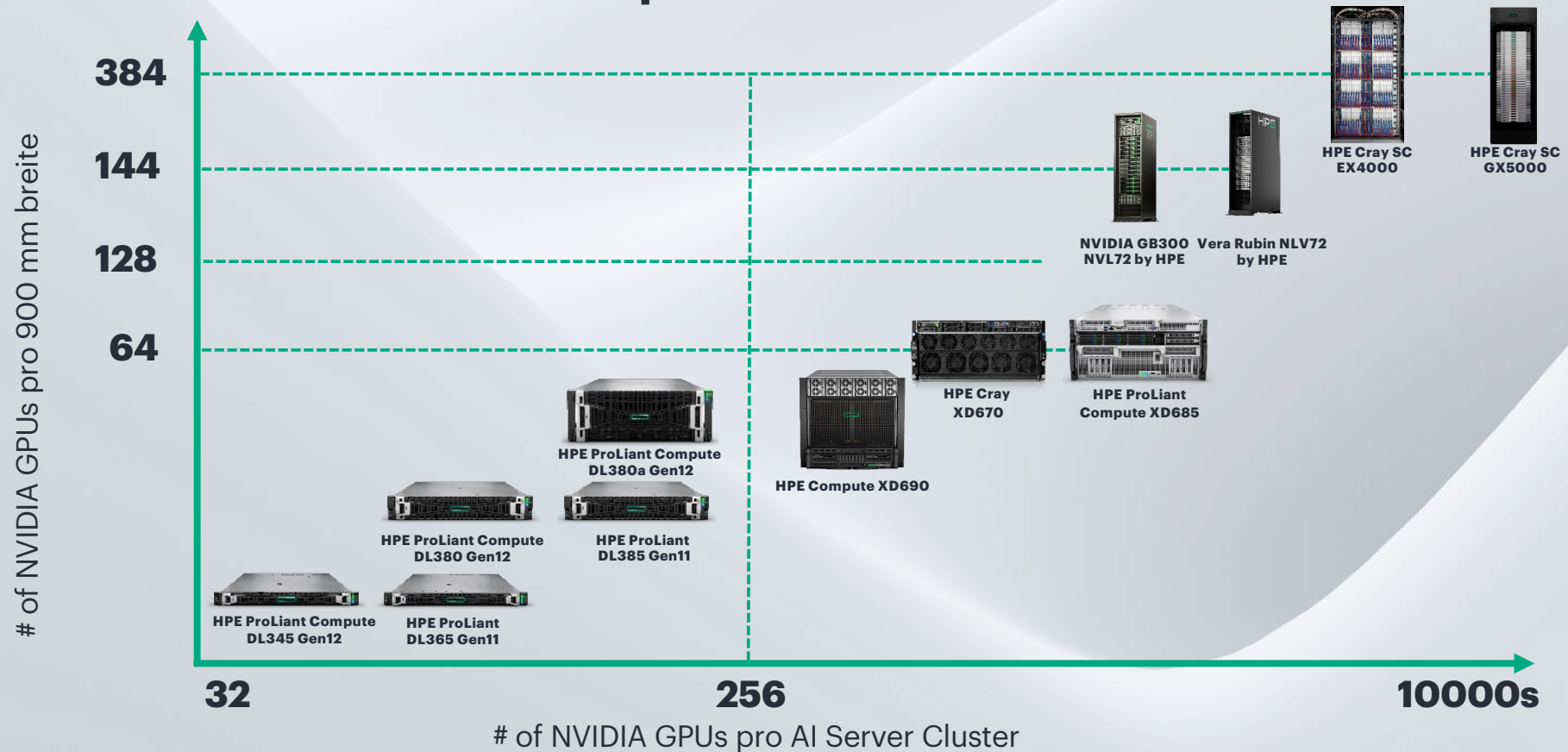
41%

48%

Energiesparnis

25%

Nicht nur das breiteste Portfolio am Markt, sondern auch die größte NVIDIA GPU Dichte pro Rack!



Performancedichte ist gut aber, kann das RZ mithalten?



„In mein Rack gehen nur 10 kW“

**„Meine Kühlung ist darauf
nicht ausgelegt“**

„Mein RZ ist so ineffizient“

„Mein RZ kann keine Wasserkühlung“

In 3 Leveln zur AI-Mastery

3.

AI Mastery durch die Nutzung von AI Plattformen:

Zentralisierung der AI Kompetenzen durch den Aufbau einer holistischen AI Plattform wie **Private Cloud AI** und dem **AI Factory Portfolio**.

2.

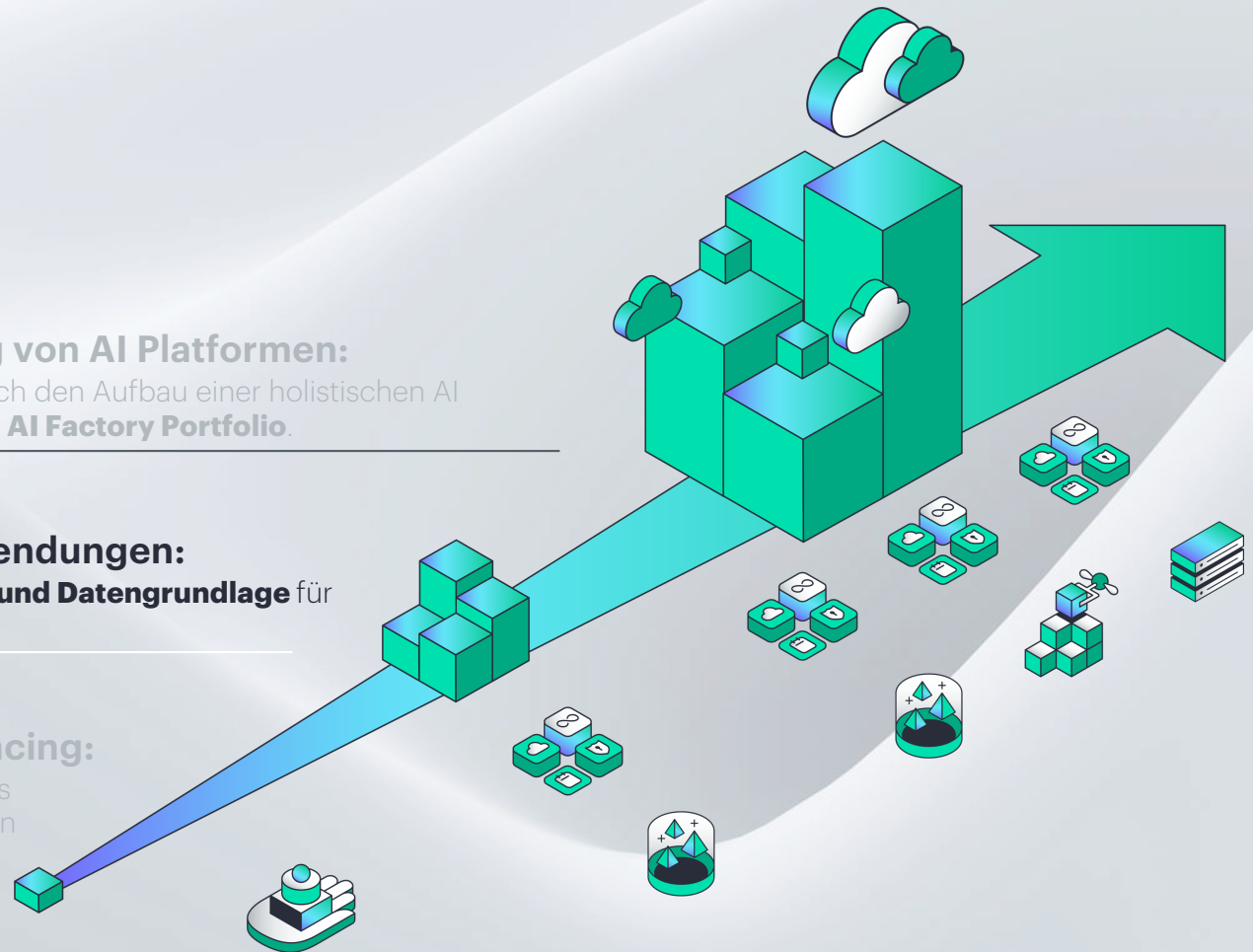
Netzwerk & Daten für AI Anwendungen:

Schaffung einer **AI-fähigen Netzwerk und Datengrundlage** für zukünftige AI Anwendungen.

1.

Rechenzentrum für AI Inferencing:

Nutzung des KI-fähigen Server Portfolios inkl. **3rd party SW** und Anforderungen an Rechenzentren.





Official Team Partner

HPE



AMG
PETRONAS
FORMULA ONE TEAM

Daten speichern reicht nicht, sie brauchen Kontext



In 3 Leveln zur AI-Mastery

3.

AI Mastery durch die Nutzung von AI Plattformen:

Zentralisierung der AI Kompetenzen durch den Aufbau einer holistischen AI Plattform wie **Private Cloud AI** und dem **AI Factory Portfolio**.

2.

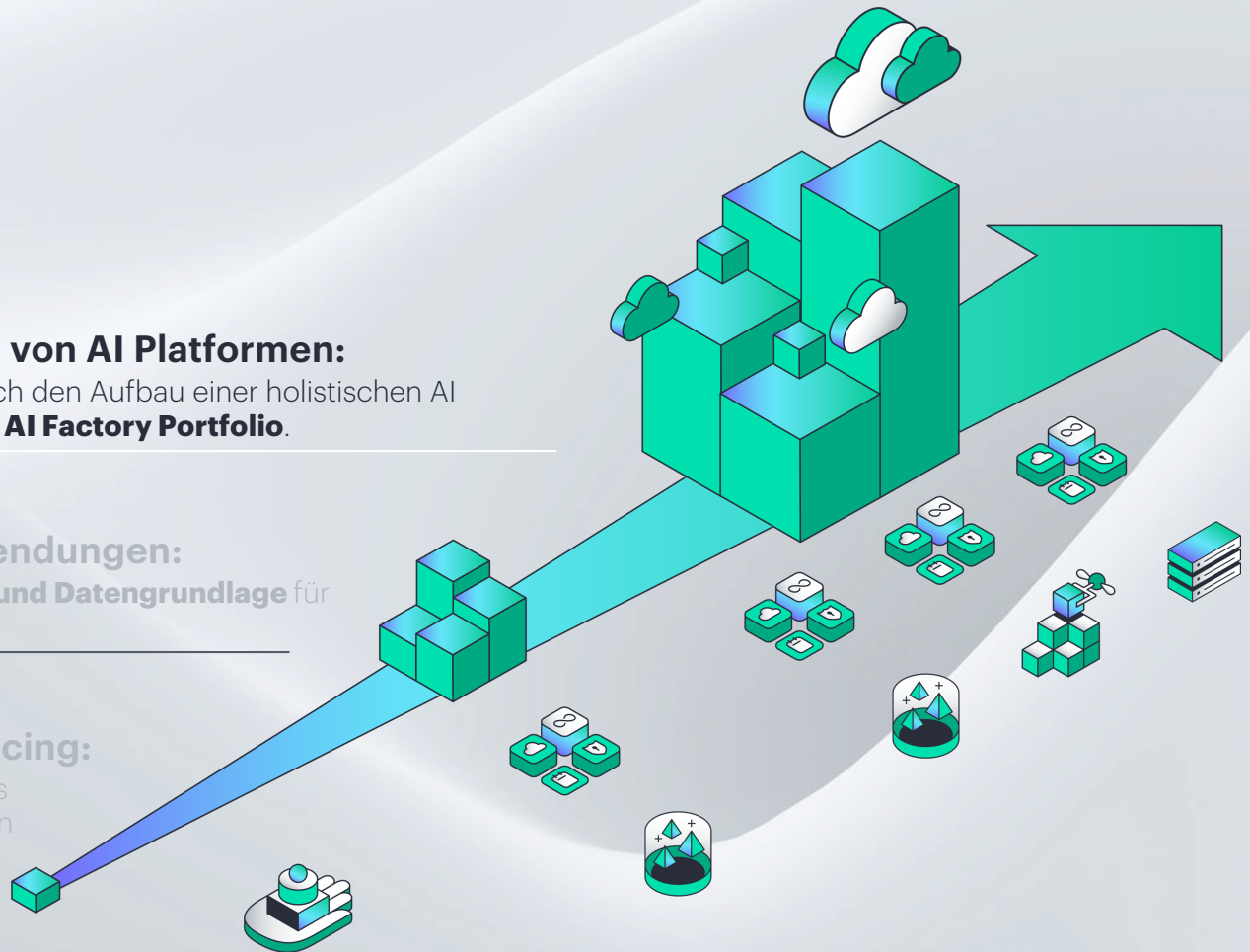
Netzwerk & Daten für AI Anwendungen:

Schaffung einer **AI-fähigen Netzwerk und Datengrundlage** für zukünftige AI Anwendungen.

1.

Rechenzentrum für AI Inferencing:

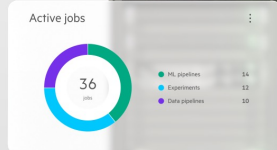
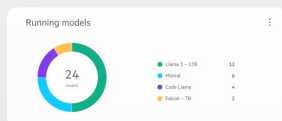
Nutzung des KI-fähigen Server Portfolios inkl. **3rd party SW** und Anforderungen an Rechenzentren.



NVIDIA AI Computing by HPE

HPE Private Cloud AI

HPE

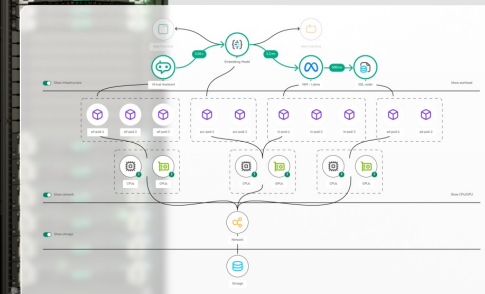
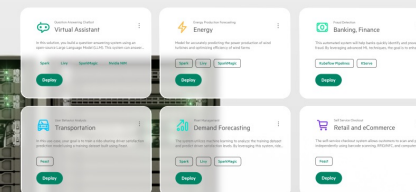


NVIDIA AI Computing by HPE

HPE Private Cloud AI

Start running AI Workloads on your AI Systems. Speed time to value for generative AI with a full-stack AI-native tuning and inference solution purpose built for the enterprise.

Launch HPE Private Cloud AI



NVIDIA AI Computing HPE Private Cloud

HPE



NVIDIA AI Computing
HPE Private Cloud AI

Start running AI Workloads on your
generative AI with a full-stack AI-native
purpose built for the enterprise.

Launch HPE Private Cloud



Open-Source Tools

Geschwindigkeit

Self-service

Zwei Welten arbeiten zusammen

Compliance

Kosten

Betrieb

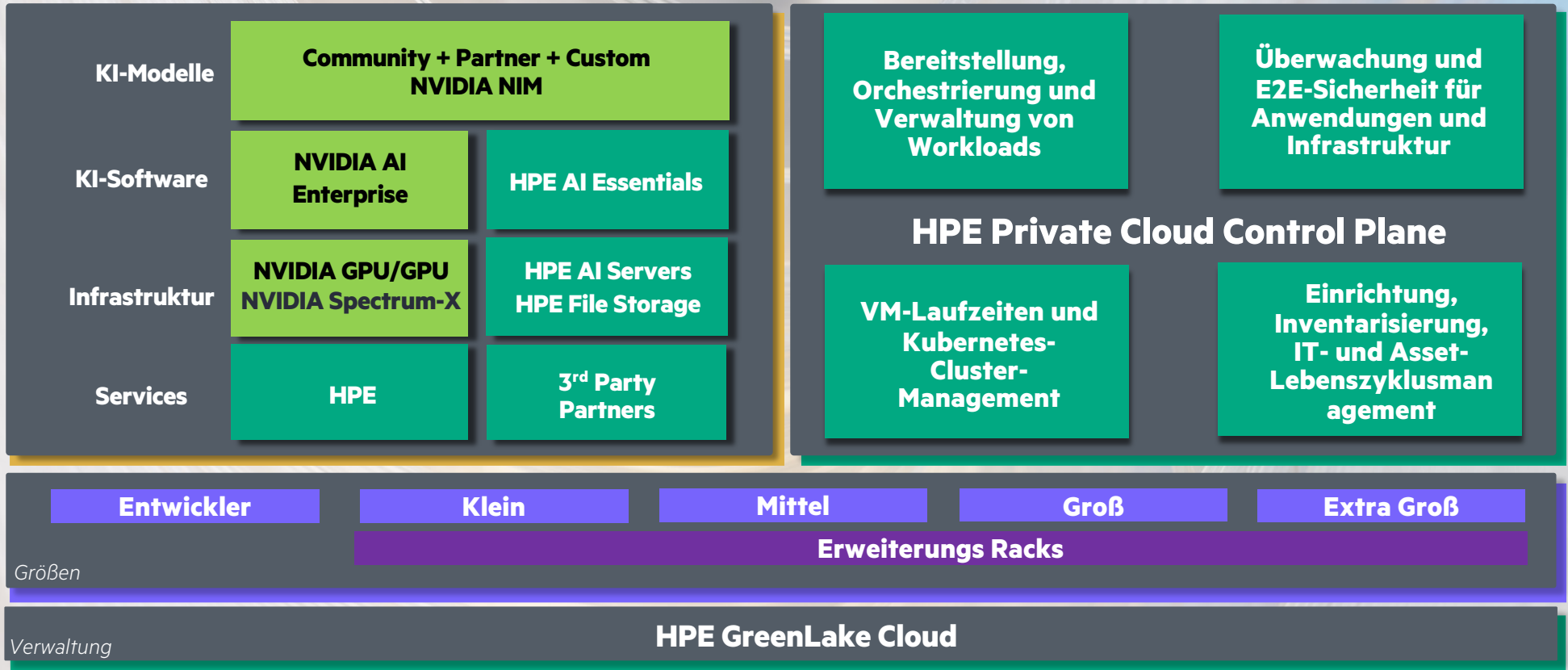


NVIDIA AI Computing by HPE

HPE Private Cloud AI

Bestandteile

Administration

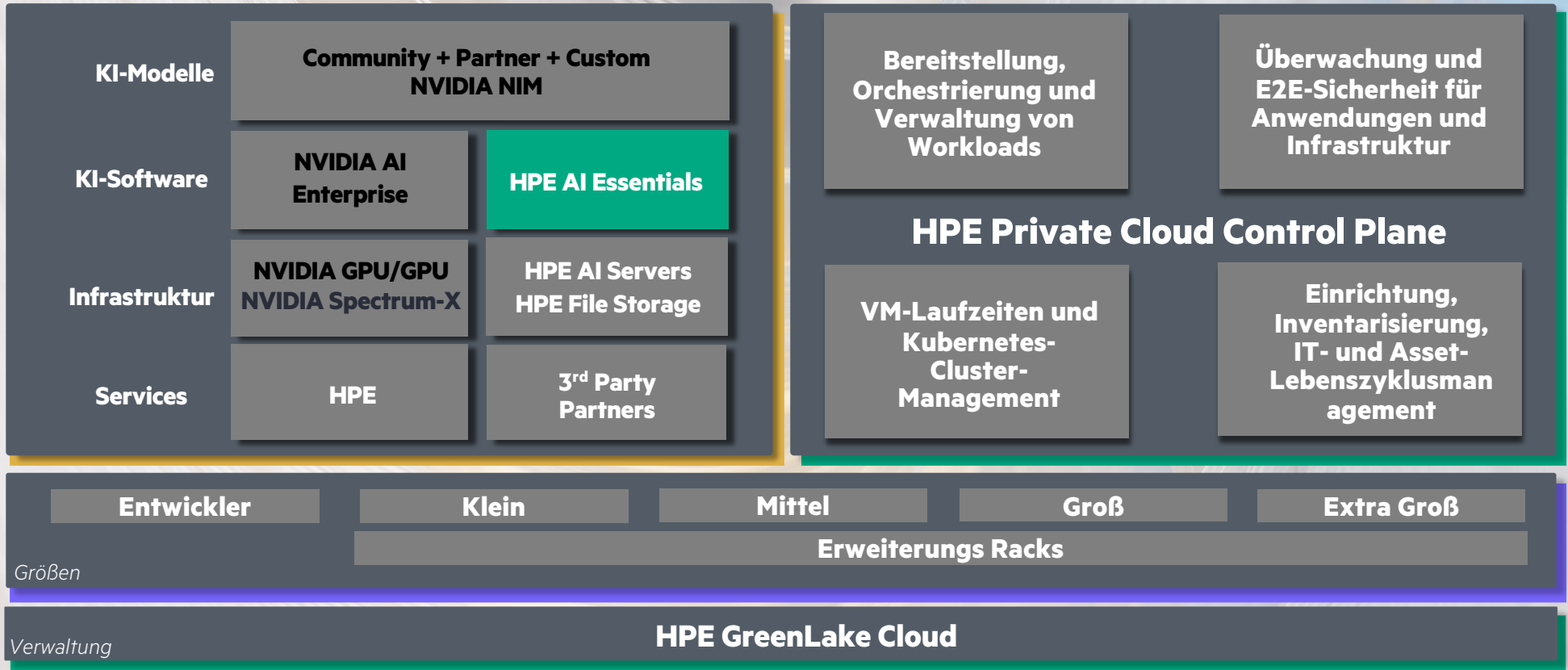


NVIDIA AI Computing by HPE

HPE Private Cloud AI

Bestandteile

Administration



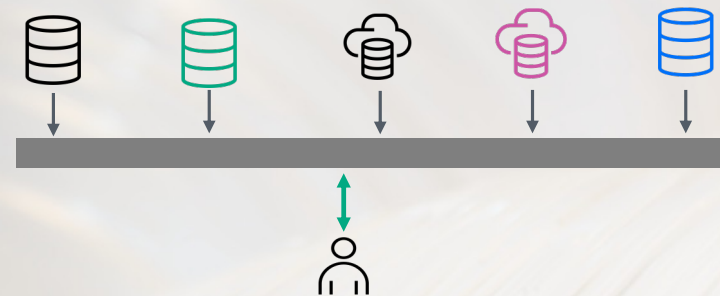
HPE AI Essentials

Kernkomponenten

Data Engineering Tools

Airflow | Presto | Data connectors

Daten aus
verschiedenen Quellen
sammeln, aufbereiten
und **bereitstellen**



Analytics & Visualization Tools

Livy | Superset | Apache Spark

Datenmengen
analysieren und
interaktive **Diagramme**,
Dashboards oder
Reports anzeigen



Data Science Tools

Kubeflow | Mlflow | Feast | Ray

Optimieren, trainieren
und **versionieren** von
KI-Modellen

LLaMA
by Meta



HPE

NVIDIA

Confidential | For Training Purposes Only

Danke!

Jeremy.massow@hpe.com

