



رویکرد مناسب ایران در

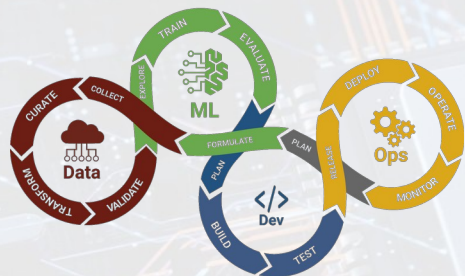
زیرساخت‌های پردازشی هوش مصنوعی

سید محمد محمدزاده ضیابری

رئیس کمیسیون هوش مصنوعی

سازمان نظام صنفی رایانه‌ای

ارکان زیرساختی هوش مصنوعی



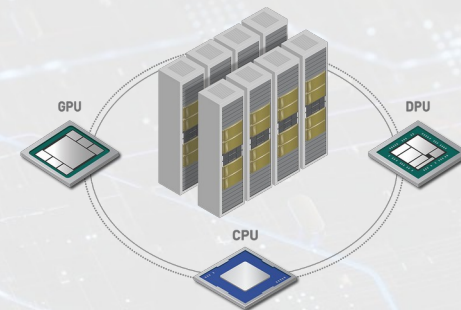
زیرساخت مهندسی هوش



زیرساخت یادگیری ماشین



زیرساخت داده‌ای



زیرساخت پردازشی



۱۹۰۰

اختراع لامپ خلاء

۱۹۲۰

اولین ربات جهان

۱۹۳۸

تولد منطق دیجیتال

۱۹۴۸

تولد علم سایبرنتیک



John McCarthy
1927-2011

۱۹۵۶

تولد کلمه
هوش مصنوعی

جان مک کارتی در مقاله‌ای
با عنوان پروپوزال دارتموث
برای اولین بار کلمه
«هوش مصنوعی»
را پیشنهاد کرد



۱۹۸۰

ربات نوازنده Wabot2

۱۹۸۲

شبکه‌های هاپفیلد

۱۹۸۵

خلق پردازش موازی

۱۹۸۹

شبکه‌های کانولوشن



۱۹۹۷

شکست کاسپاروف

ابراهیم دیپ‌بلو از شرکت
IBM توانست ابرشرط‌نچ‌باز
جهان را در مسابقه شطرنج
شکست دهد!



۲۰۰۳-۲۰۰۵

مسابقات دارپا

۲۰۰۶

شبکه‌های باور عمیق

۲۰۰۸

مغز گوگل

۲۰۱۱

معرفی سیری اپل



۲۰۱۷

شبکه‌های ترنسفرمر

"
Attention
Is
All
You
Need
"



۲۰۱۷

شبکه‌های wavenet

۲۰۱۸

GPT-1

۲۰۱۹

مدل زبانی BERT

۲۰۲۰

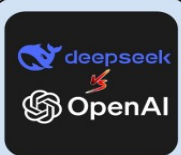
GPT-3



۲۰۲۲

ChatGPT

چت‌جی‌بی‌تی
تنها در طول دو ماه
توانست بیشترین میزان
محبوبیت و نصب را در
میان همه نرم‌افزارها
کسب کند!



۲۰۲۵

deepseek

چت بات چینی دیپ‌سیک
تنها در دو هفته توانست
تعداد نصب چت‌جی‌بی‌تی
را پشت سر گذاشته و به پر
استفاده‌ترین نرم‌افزار هوش
مصنوعی بدل شود!

خزان اول

خزان دوم

۱۹۰۰ - ۱۹۵۰

۱۹۵۰ - ۱۹۷۳

۱۹۸۰ - ۱۹۸۷

۱۹۹۳ - ۲۰۰۰

۲۰۰۰ - ۲۰۱۲

۲۰۱۲ - ۲۰۱۷

۲۰۱۷ - ۲۰۲۱

۲۰۲۱ - ۲۰۲۵

۲۰۲۵ - ...

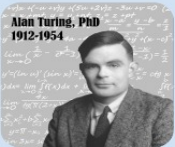
۲۰۱۰



۱۹۵۰

تست تورینگ

ماشین زمانی شایسته
هوشمند نامیده شدن
است که بتواند انسان را به
نحوی فریب دهد که
انسان فرض شود



Alan Turing, PhD
1912-1954

۱۹۵۷

اولین برنامه هوشمند

۱۹۵۸

شبکه پرسپترون

۱۹۶۱

اولین ربات صنعتی

۱۹۶۶

چت بات الیزا



Frank Rosenblatt
1928 - 1971

۱۹۸۶

پس انتشار خطا

شبکه‌های چندلایه با قابلیت
پس انتشار خطا
انقلابی در قابلیت‌های
شبکه‌های عصبی مصنوعی
ایجاد کردند



۱۹۹۶

تولد رباتکاپ

۱۹۹۷

شبکه‌های LSTM

۱۹۹۹

تولد سگ‌های آیبو

۲۰۰۰

تولد ربات آسیمو



۲۰۱۲

AlexNet

الکسنت شبکه عصبی
مصنوعی ۸ لایه با ۵ لایه
کانولوشن توانست شبکه
عصبی را ImageNet
شکست دهد و از GPU
برای آموزش شبکه
استفاده شد.



۲۰۱۳

DeepMind

۲۰۱۴

تولد شبکه‌های مولد

۲۰۱۵

Tensorflow

۲۰۱۶

Alpha Go



AlphaGo
Lee Sedol

۲۰۲۱

DALL-E

خلق تصویر از روی
توضیحات، قدرت هوش
مصنوعی را در انجام
وظایف چند وجهی
به نمایش می‌گذارد



۲۰۲۰

خودرو خودران تسلا

۲۰۲۱

قاعده‌گذاری و اخلاق

۲۰۲۴

مدل تولید ویدیو sora

۲۰۲۴

EU AI Act



+۲۰۲۵

AGI

ASI

تکینگی

ابرانسان‌ها

و

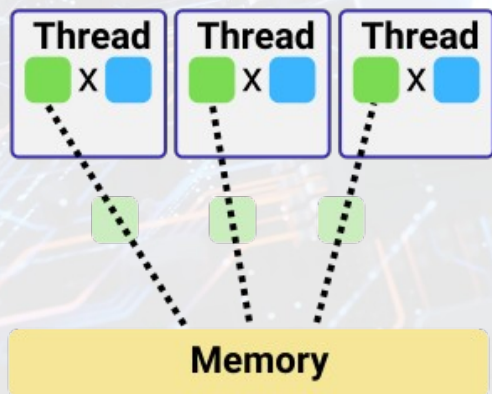
فرا تر از آن



انواع پردازنده‌های هوش مصنوعی

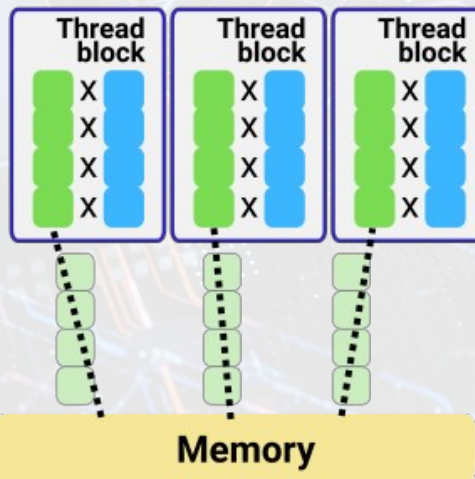


CPU



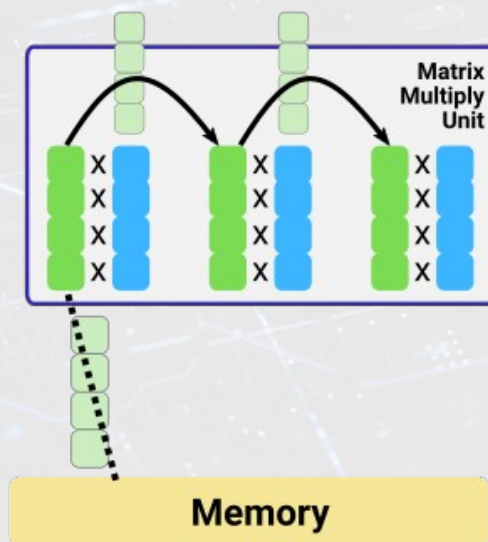
ضرب موازی خطی

GPU



ضرب موازی برداری

TPU



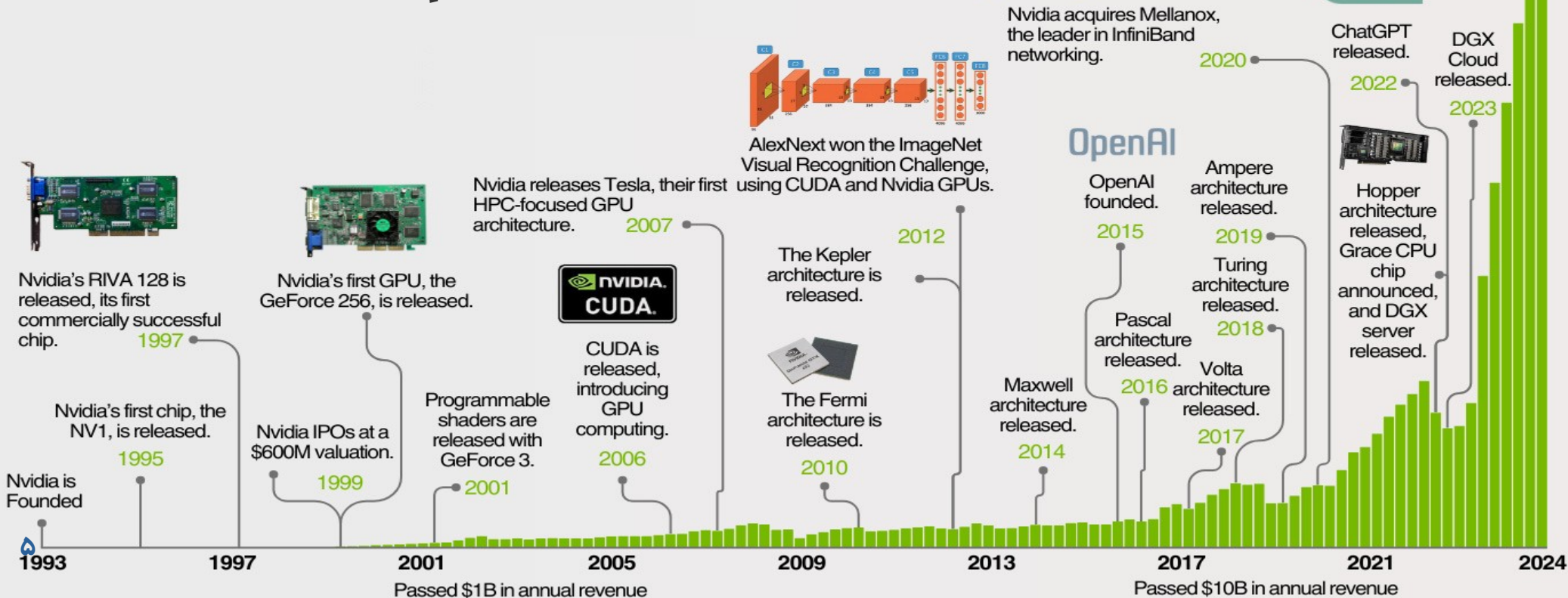
ضرب موازی ماتریسی



Nvidia surpasses \$35B in quarterly revenue, up 5x in just 6 quarters.

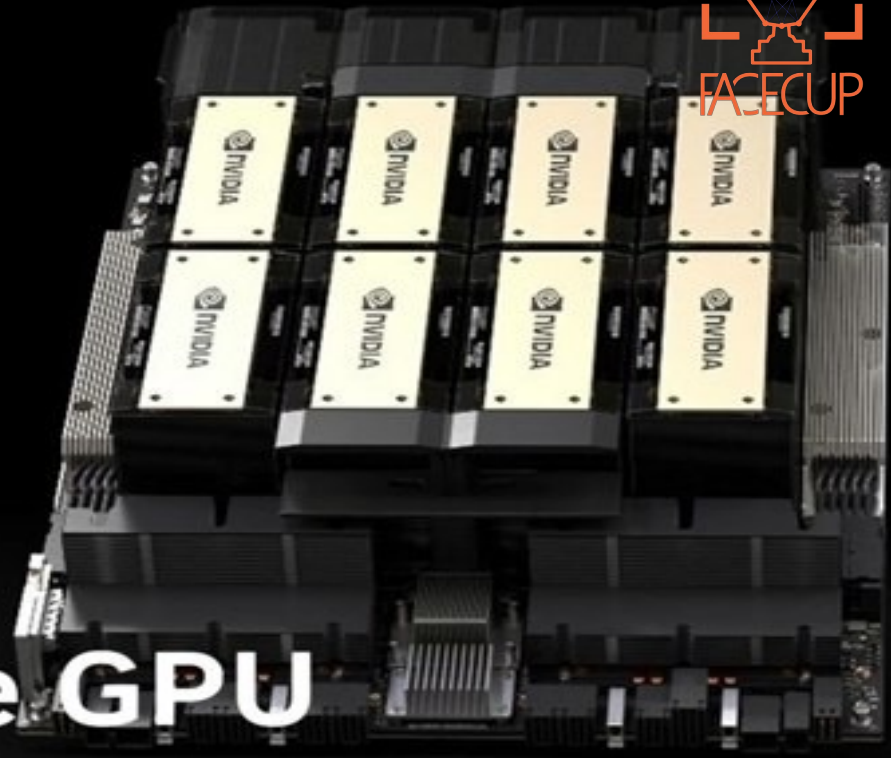
2024

تغییرات سهام Nvidia



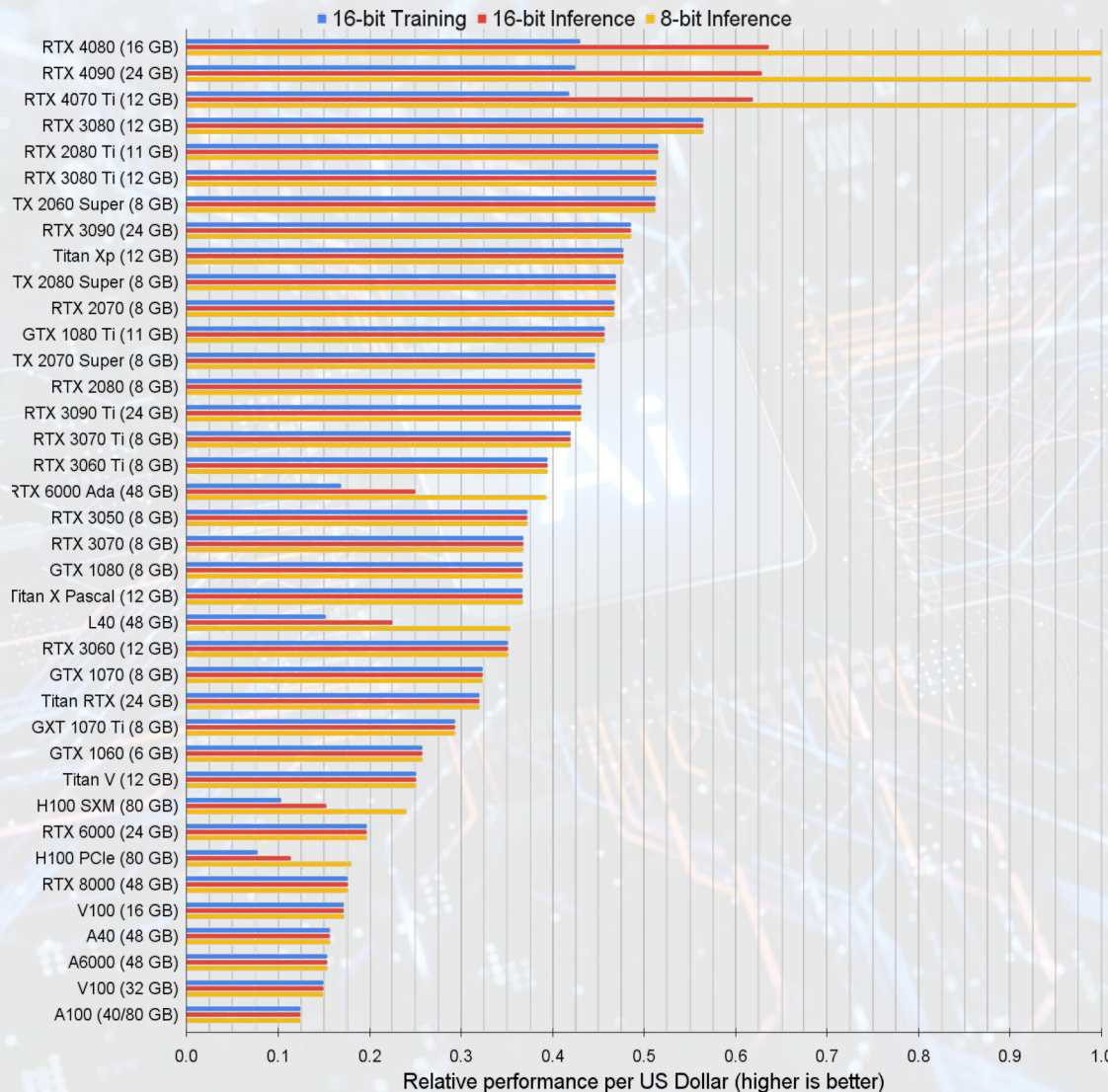


H200
Tensor Core GPU



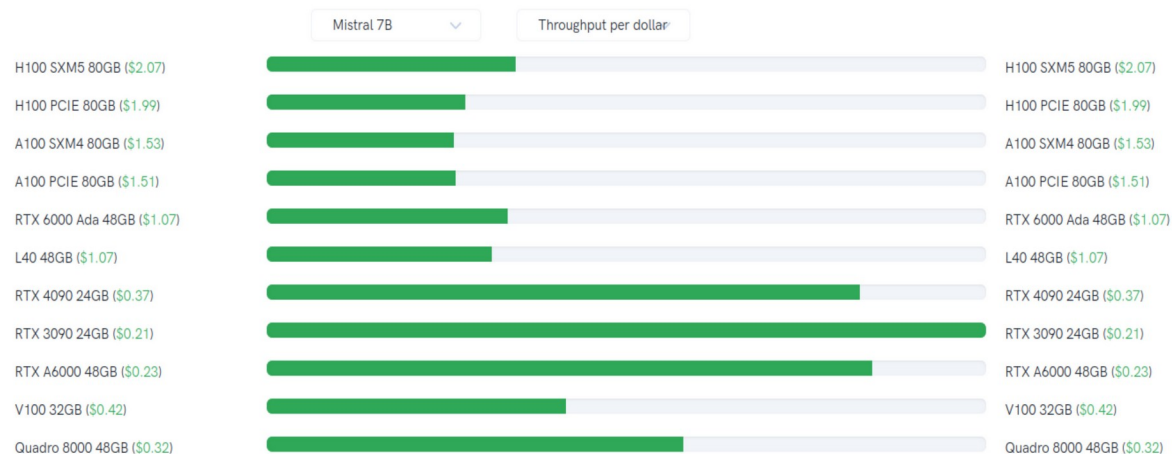


کارایی به قیمت کارت‌های گرافیک در پردازش‌های ۱۶ بیت



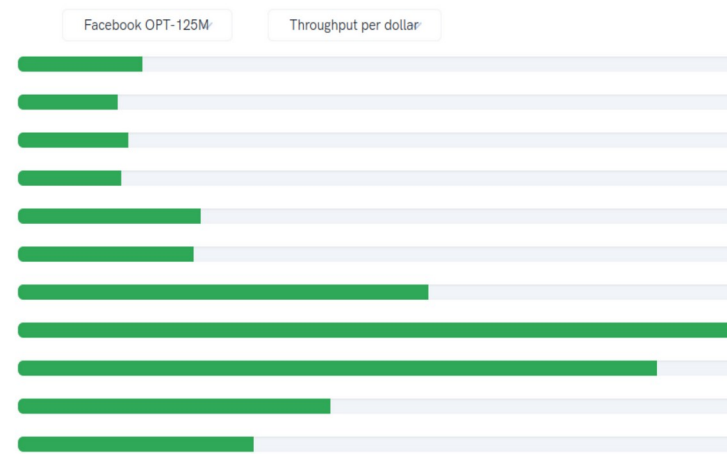
LLM Inference Throughput

Relative tokens per second on Mistral 7B. Half precision (FP16). The more, the better.



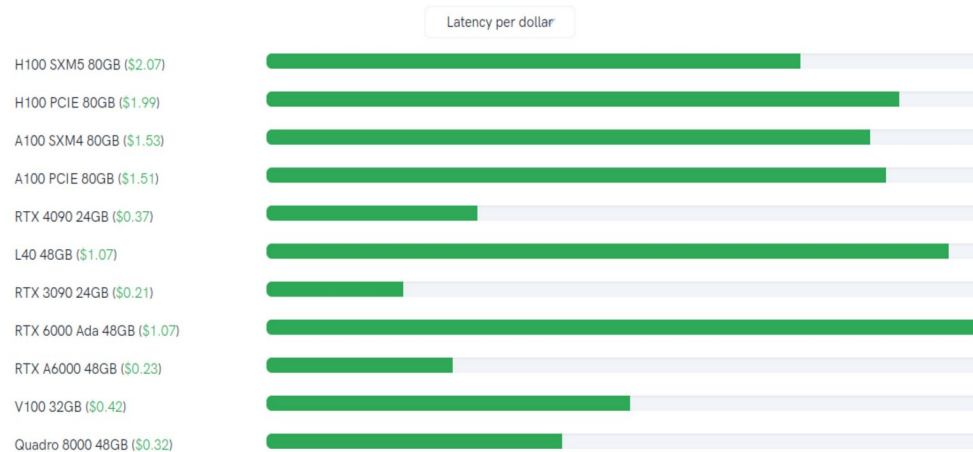
LLM Inference Throughput

Relative tokens per second on Mistral 7B. Half precision (FP16). The more, the better.



LLM Batch Latency

Time taken to process one batch of tokens, p90, Mistral 7B. Half precision (FP16). The lower, the better.



PyTorch Training GPU Benchmarks

Visualization

chart

Metric

throughput/\$

Precision

fp16

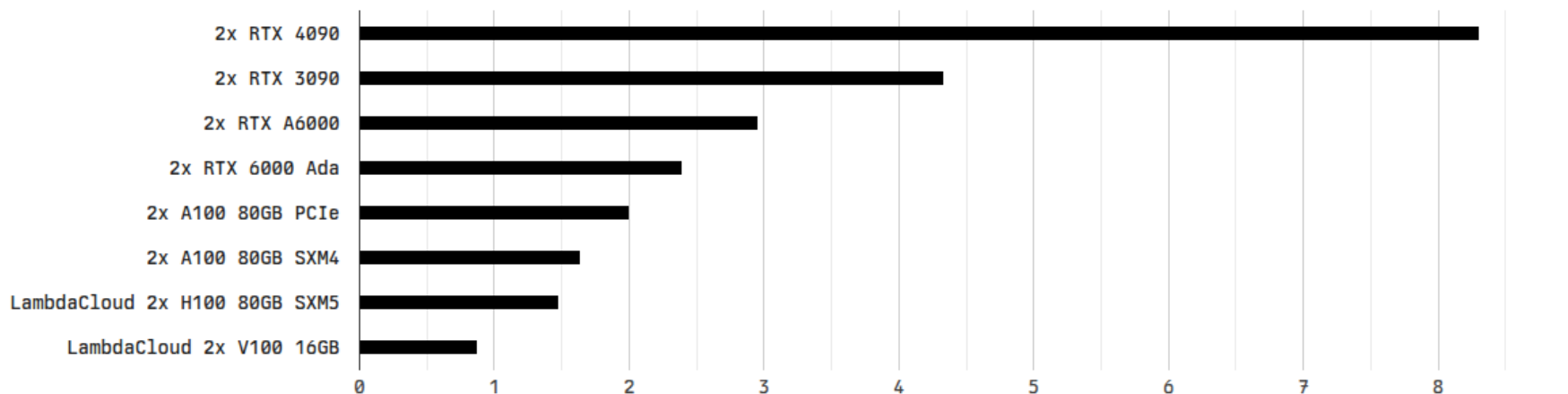
Number of GPUs

2x

Model

All Models

Relative Training Throughput Per \$ w.r.t 1xLambdaCloud 1x V100 16GB (All Models)



* Higher is better

<https://lambda.ai/gpu-benchmarks>

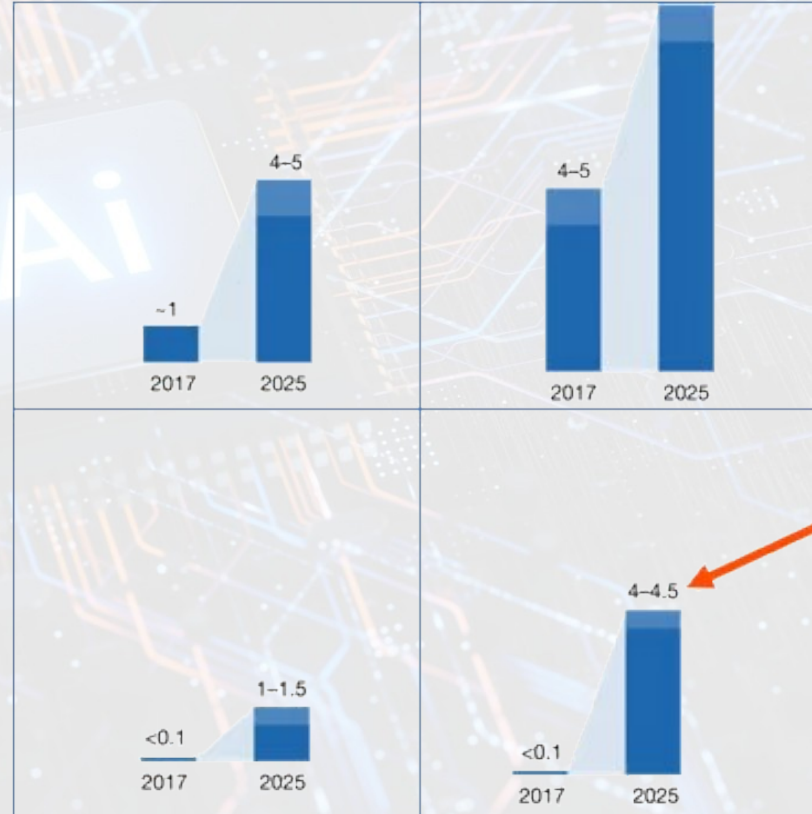
سرمایه گذاری ها از سال ۲۰۱۷ تا کنون



Data Center

Training

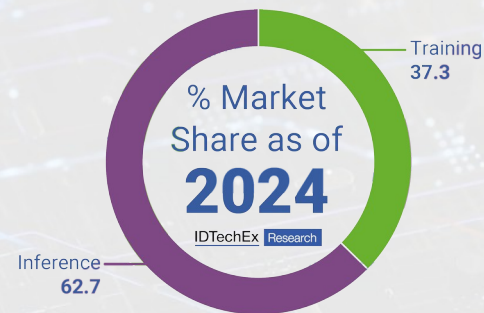
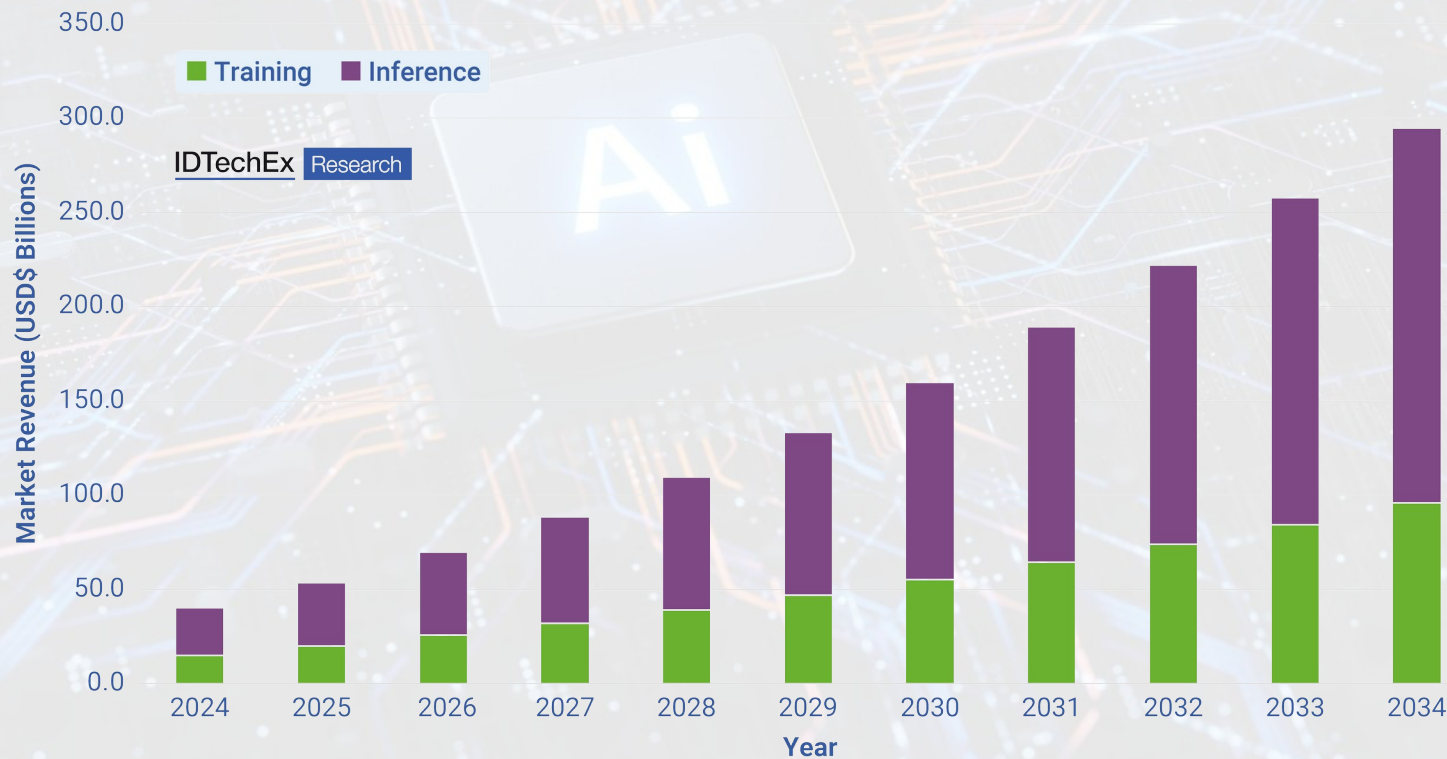
Inference



Billions!



پیش بینی تغییرات بازار در دو حوزه آموزش و استنتاج



کارت متناسب با کاربرد



 **DGX A100**

 **DGX H100**

 **DGX H200**

 **Instinct Mi300**

 **Ascend 910B**

Inference

Training

 **GEFORCE RTX 2080**

 **GEFORCE RTX 3060**

 **GEFORCE RTX 3090**

 **GEFORCE RTX 4060**

 **GEFORCE RTX 5060**

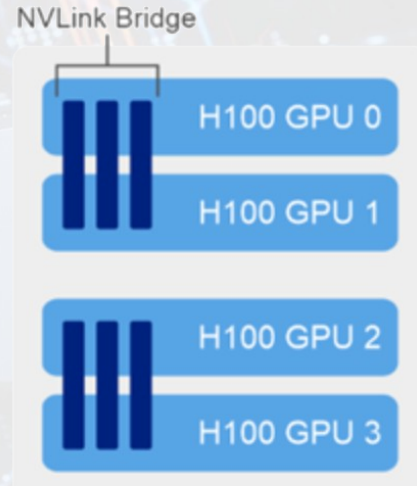
 **GEFORCE RTX 4090**

 **GEFORCE RTX 5090**

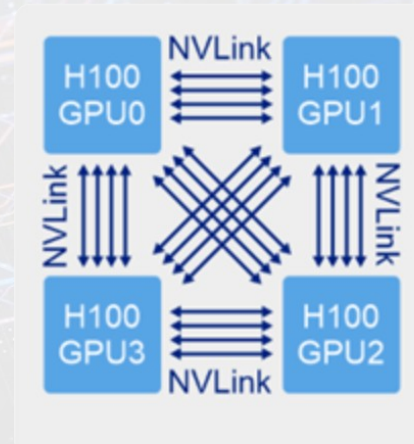
انواع اتصالات کارت‌ها



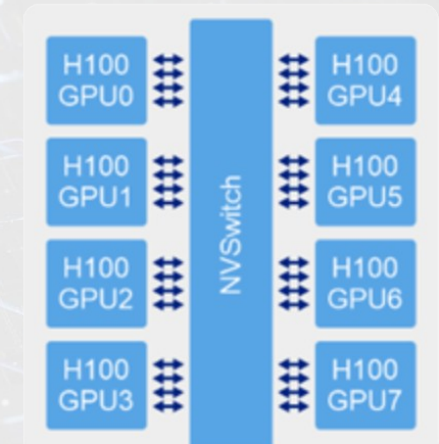
کارت‌های مستقل
PCI



زوج کارت PCI
متصل با NVLink



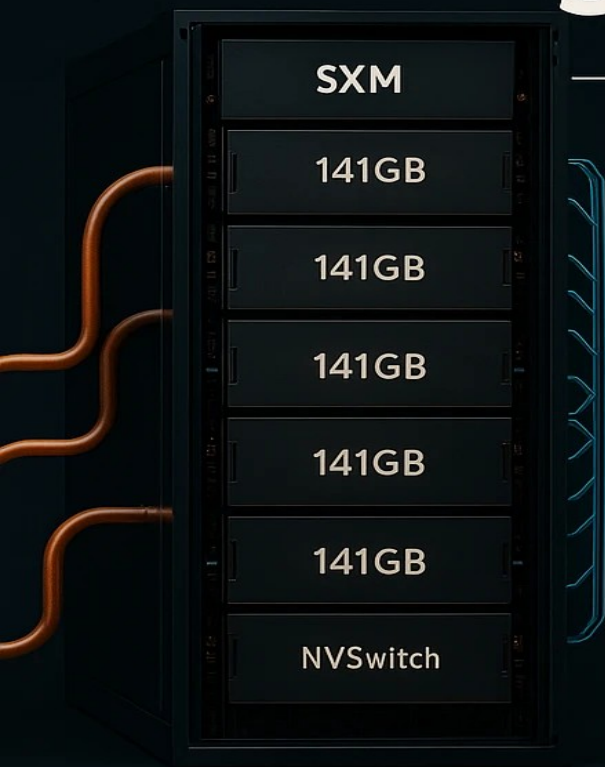
اتصال Mesh بین
کارت‌های SXM



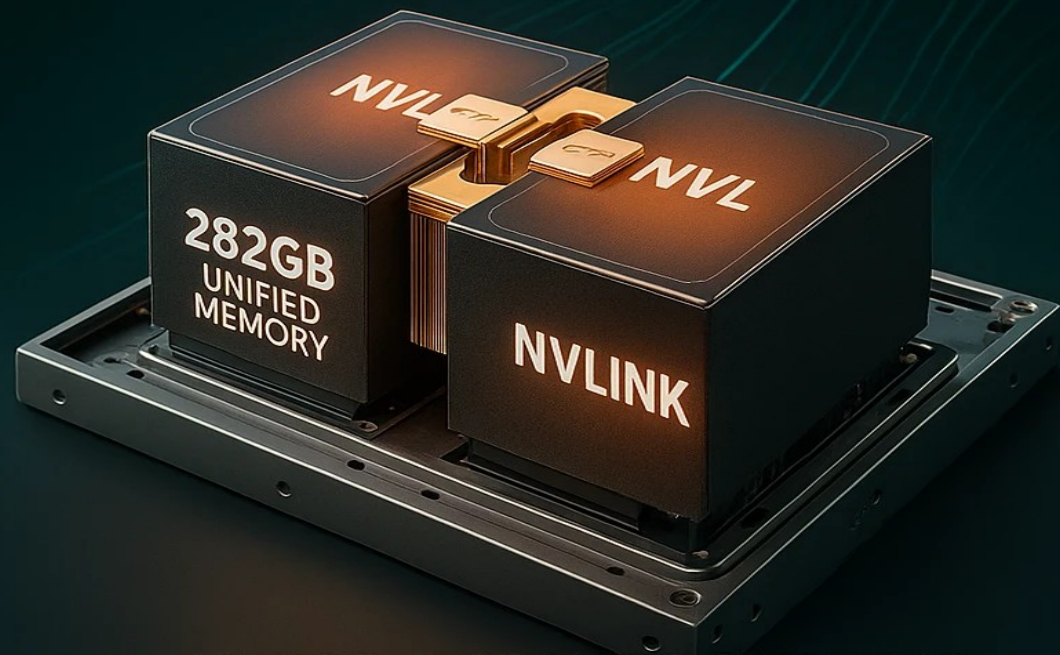
اتصال بین سرورها
با NVSwitch

SXM vs NVL

— DESIGN MATTERS —



HORIZONTAL GPU DENSITY
H200



VERTICAL SCALING SIMPLIFIED
H200 NVL

مقایسه SXM با NVL



Feature	H200 SXM	H200 NVL
Total Memory	141GB per GPU	282GB unified pool
Memory Bandwidth	4.8 TB/s per GPU	9.6 TB/s total
NVLink Speed	900 GB/s (between GPUs)	1.8 TB/s (internal bridge)
Ideal For	Models fitting within 141GB per GPU	AI models >141GB
Llama 70B, GPT-4 inference	Requires sharding across GPUs Communication delays between GPUs	1.6–2× throughput Entire model fits in 282GB unified memory
Image upscaling, Financial modeling	Matches NVL per-GPU performance	No advantage (Same ~67 TFLOPS FP64/GPU)

انواع سرورهای Nvidia



DGX™

HGX™

MGX™

EGX™



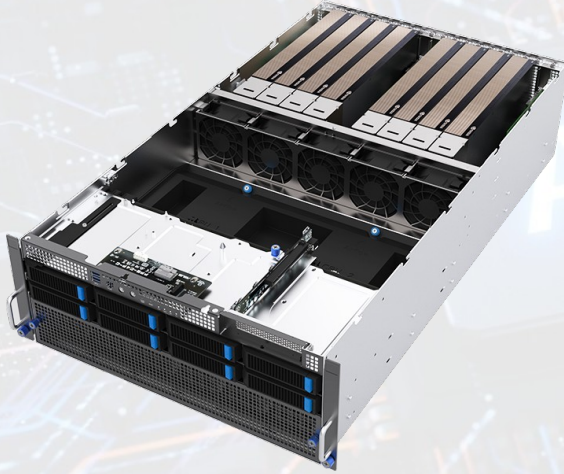
سرور آماده همراه
نرم افزار هوش مصنوعی
SXM

سرور منعطف برای
کاربردهایی مختلف
SXM

سیستم های مناسب
برای پردازش لبه
PCIe

سکوی مرجع ماژولار و
انعطاف پذیر
PCIe

انواع دیگر سرورها



SUPERMICRO®



ASUS®

GIGABYTE



مقایسه سرورهای PCI با هم

سازنده	بیشترین تعداد کارت	پشتیبانی از RTX4090	مهمترین مزایا	مهمترین معایب
 NVIDIA	۸	خیر	عملکرد بالا با NVLink، بهینه‌سازی کامل برای AI/HPC، مقیاس‌پذیری آسان در کلاسترها	گران‌قیمت (از ۲۰۰,۰۰۰ دلار)، انعطاف‌پذیری کم برای سفارشی‌سازی
 HEWLETT PACKARD	۸	خیر	پشتیبانی enterprise قوی، امنیت بالا، ادغام با hybrid cloud	رشد کندتر در بازار AI، گران‌تر برای حجم کاری بالا، پشتیبانی محدود از کارت‌ها
 SUPERMICRO	۲۰	بله	هزینه رقابتی، انعطاف‌پذیری در PCIe و پشتیبانی از GPUهای متنوع	مصرف برق بالا در کانفیگ‌های متراکم، پشتیبانی ضعیف‌تر نسبت به enterprise
 GIGABYTE	۱۶	بله	ارتقاپذیری آسان، دارای گزینه‌های liquid cooling	بهینه‌سازی کمتر برای workloads سنگین AI نسبت به NVIDIA
 ASUS	۸	بله	خنک‌سازی عالی، PCIe 5.0 برای آینده‌نگری، کارایی بالا در inference	کمتر شناخته‌شده
RIG	۲۰	بله	آزادی عمل بالا، هزینه بسیار پایین	فاقد استانداردهای مراکز داده

انواع خدمات پردازشی

- **Dedicated:** سرورهای اختصاصی به صورت تکی و خوشه‌های پردازشی
- **VM/Cloud:** ماشین‌های مجازی یا ابری با پردازنده گرافیکی اختصاصی به صورت تکی یا خوشه‌ای
- **VPDC:** مرکز داده اختصاصی به صورت مجازی شامل تمامی امکانات پردازشی و شبکه به صورت مدیریت شده



GPUDN



Shared GPU



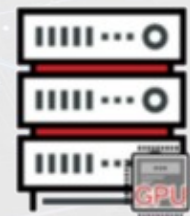
Sliced GPU



VPDC



VM/Cloud



Dedicated

- **Sliced GPU:** پردازنده گرافیکی قطعه‌بندی شده و تخصیص داده شده به صورت اختصاصی (صرفاً برای کارت‌های حرفه‌ای)
- **Shared GPU:** فناوری اشتراک‌گذاری واقعی پردازنده گرافیکی به صورت کاملاً پویا قابل بهره‌برداری بر روی تمامی کارت‌ها
- **GPUDN:** شبکه پردازش هوش مصنوعی نزدیک به مصرف‌کننده مورد استفاده در مناطق محروم و با شبکه ضعیف



جمع بندی

- مشکلات بومی راهکار بومی نیاز دارد
- هیچ گزینه حلال المسایلی وجود ندارد!
- ترندهای جهانی لزوماً برای ایران مناسب نیستند
- تفکیک دقیق و الزامی زیرساختهای استنتاج از آموزش
- با مهندسی هوش مصنوعی می توان نیاز به سخت افزار را بهینه کرد
- بهره وری باید اصل اول در انتخاب سخت افزار باشد نه ترند بودن یا مقالات

برای دریافت اسلایدها روی پیامرسانهای بله یا واتساپ به شماره ۰۹۲۰۳۳۰۲۴۶۱ پیام بدهید