

EnQ
Digital

SUPERCOMPUTAÇÃO EXASCALE

EL CAPITAN

Arquitetura, desempenho, energia, liquid cooling e infraestrutura de data center

Edição técnica | Agosto 2026

EnQ Digital | Cloud - IA - Data Center - Conectividade

VISÃO EXECUTIVA

El Capitan em 60 segundos

Uma leitura de engenharia para quem projeta infraestrutura de IA, HPC e data centers de alta densidade.

1,809 EF/s

HPL Rmax - jun/2026

2,821 EF/s

Rpeak TOP500

29,685 MW

potência HPL medida

60,94 GF/W

eficiência HPL

O El Capitan é um HPE Cray EX255a instalado no Lawrence Livermore National Laboratory (LLNL), na Califórnia. Ele foi projetado para a National Nuclear Security Administration (NNSA) e combina computação exascale, inteligência artificial, simulação multiphysics, armazenamento paralelo e liquid cooling em uma única plataforma de escala industrial.

Em novembro de 2024, o sistema estreou no TOP500 com 1,742 exaFLOP/s. Uma nova medição elevou o resultado para 1,809 exaFLOP/s em novembro de 2025. Na lista de junho de 2026, El Capitan passou ao segundo lugar global, mantendo os 1,809 exaFLOP/s e permanecendo uma das referências mais importantes de arquitetura exascale do planeta.

Ponto de engenharia

O valor de 29,685 MW do TOP500 é a potência medida durante o benchmark HPL. O LLNL publica ~35-36 MW como ordem de grandeza de pico do sistema, enquanto a modernização ECFM elevou a capacidade elétrica do complexo para 85 MW. Esses três números representam camadas diferentes da infraestrutura.

SUMÁRIO

O que este e-book cobre

Bloco	Conteúdo
01	Exascale, benchmarks e posicionamento global
02	AMD MI300A e arquitetura do compute node
03	Escala física: nós, APUs, memória e gabinetes
04	HPE Slingshot 11 e topologia Dragonfly
05	Rabbit near-node storage e Lustre
06	TOSS 4, Flux, ROCm e software stack
07	Potência elétrica, energia anual e densidade por rack
08	ECFM: 115 kV, subestação, 13,8 kV e 85 MW
09	Direct liquid cooling e planta térmica
10	Lições para data centers de IA e GPU no Brasil

Escopo público

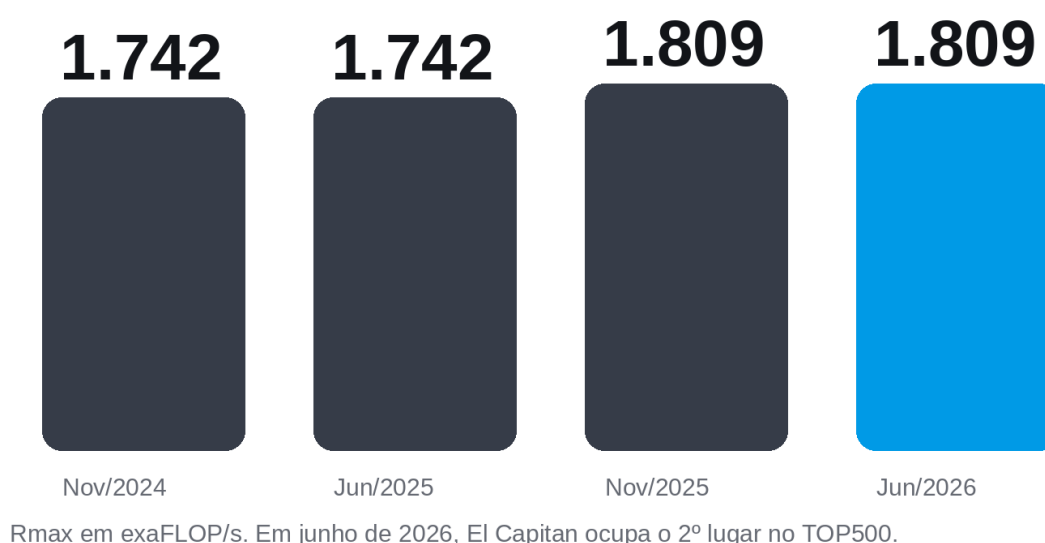
El Capitan atende missões classificadas. Este material usa exclusivamente informações técnicas públicas. Diagramas de energia e cooling representam a arquitetura funcional divulgada; não reproduzem detalhes de segurança, redundância ou proteção que não sejam públicos.

CAPÍTULO 01

O que significa operar em exascale?

Exascale é a classe de sistemas capazes de executar pelo menos 10^{18} operações de ponto flutuante por segundo em workloads de referência. Essa fronteira muda a escala de problemas científicos: mais resolução espacial, mais variáveis, mais cenários, melhor quantificação de incerteza e maior uso de técnicas de inteligência artificial integradas à simulação.

Desempenho HPL: a evolução da medição



Evolução do HPL medido do El Capitan. Fonte: TOP500 e LLNL.

O HPL mede desempenho de álgebra linear densa e continua sendo a métrica principal do TOP500. O HPCG procura refletir padrões de acesso a memória e comunicação mais próximos de aplicações científicas. Em junho de 2026, El Capitan registrava 17,406 PF/s no HPCG. No HPL-MxP, que explora precisão mista e aproxima workloads de IA, o sistema alcançou 16,7 exaFLOP/s em uma medição publicada pelo TOP500.

CAPÍTULO 02

Arquitetura geral: HPE Cray EX255a + AMD MI300A



El Capitan no LLNL. Foto: Garry McLeod/LLNL.

A arquitetura foi desenhada como um sistema integrado. Compute, fabric, storage e cooling não são subsistemas independentes; o desempenho final depende do equilíbrio entre todos eles. O HPE Cray EX foi concebido para cargas de alta densidade, com refrigeração líquida direta e rede Slingshot integrada.

90

gabinetes compute

11.520

nós totais

46.080

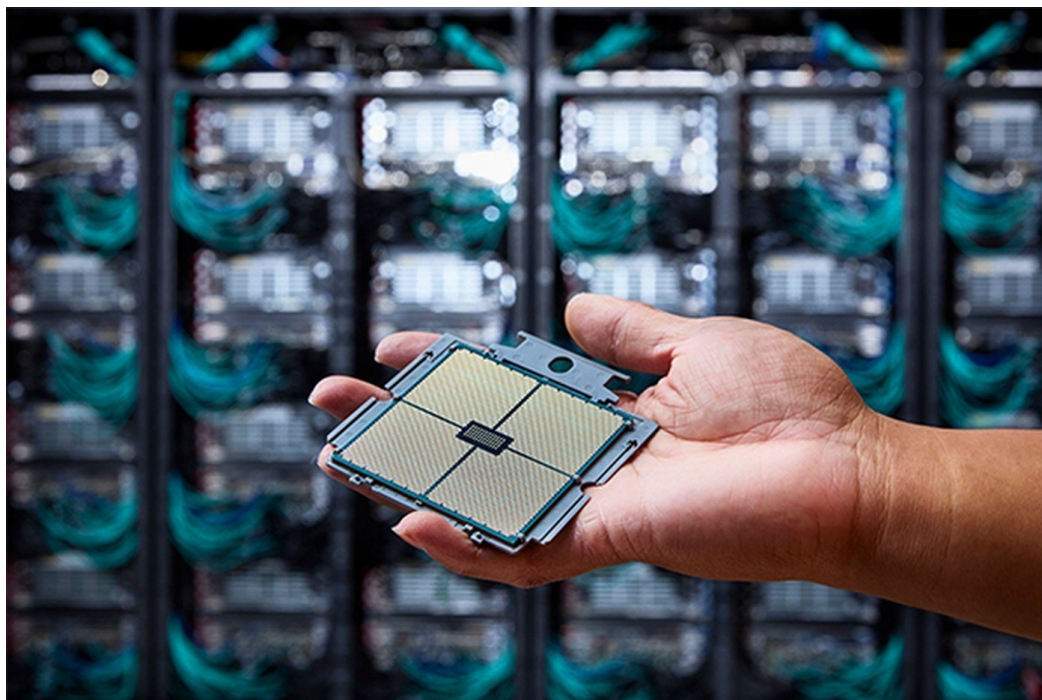
MI300A APUs

720

Rabbit modules

CAPÍTULO 03

AMD Instinct MI300A: CPU e GPU no mesmo pacote



AMD Instinct MI300A em ambiente do El Capitan. Foto: LLNL.

O MI300A é uma APU para HPC: integra CPU Zen 4, GPU CDNA 3 e memória HBM3 compartilhada no mesmo pacote. Essa memória unificada reduz a necessidade de copiar dados entre memórias fisicamente separadas de CPU e GPU - uma operação que, em arquiteturas discretas, adiciona latência, consumo e complexidade de programação.

Componente	Por MI300A
CPU	24 cores Zen 4 x86
GPU	6 XCDs CDNA 3
Memória	128 GiB HBM3
Bandwidth HBM3	~5,3 TB/s teóricos
FP64 vetorial	61,3 TFLOP/s
Coerência	CPU/GPU com memória compartilhada e Infinity Fabric

Anatomia do compute node

Nó de computação: 4 AMD Instinct MI300A

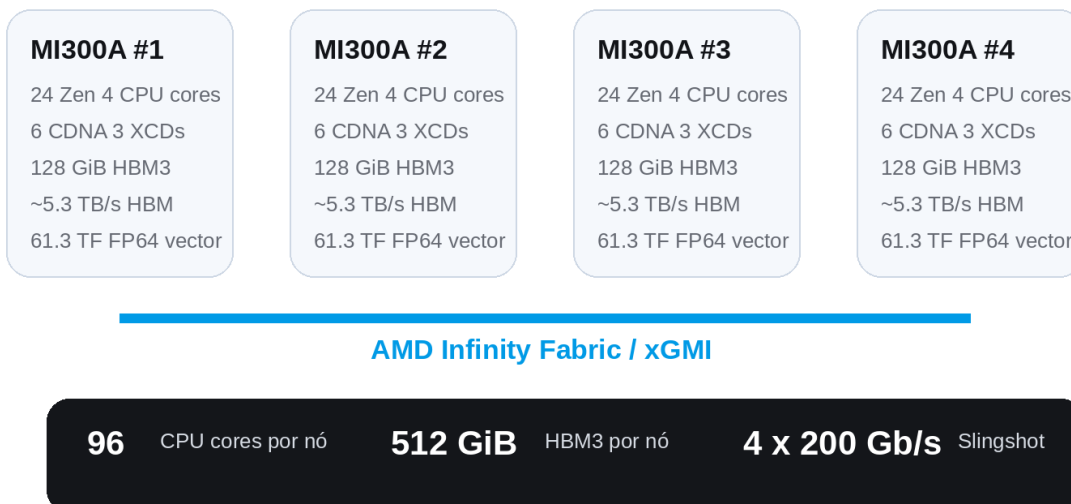


Diagrama EnQ Digital baseado na documentação pública do HPC @ LLNL.

Cada compute node possui quatro MI300A. Isso significa 96 cores de CPU Zen 4, quatro complexos acelerados com memória HBM3 e 512 GiB de memória global no nó. O próprio LLNL publica 250,8 TFLOP/s de pico FP64 por nó.

A conexão externa do nó usa quatro interfaces HPE Slingshot de 200 Gb/s. Em termos agregados, a injeção de rede chega a 100 GB/s por nó. Em uma máquina com mais de onze mil nós, o fabric deixa de ser uma rede auxiliar e passa a ser parte do computador.

Por que isso importa para IA?

Treinamento distribuído e simulação acelerada são limitados pelo tempo de movimento de dados. Uma arquitetura unificada reduz tráfego interno CPU-GPU; um fabric de alta largura de banda reduz a espera entre nós.

Escala: 11.520 nós e 46.080 APU's

Escala física do sistema



Fonte: página atual de hardware do HPC @ LLNL. O conjunto usado em benchmarks pode variar da configuração total instalada.

Escala instalada segundo o portal atual de hardware do LLNL.

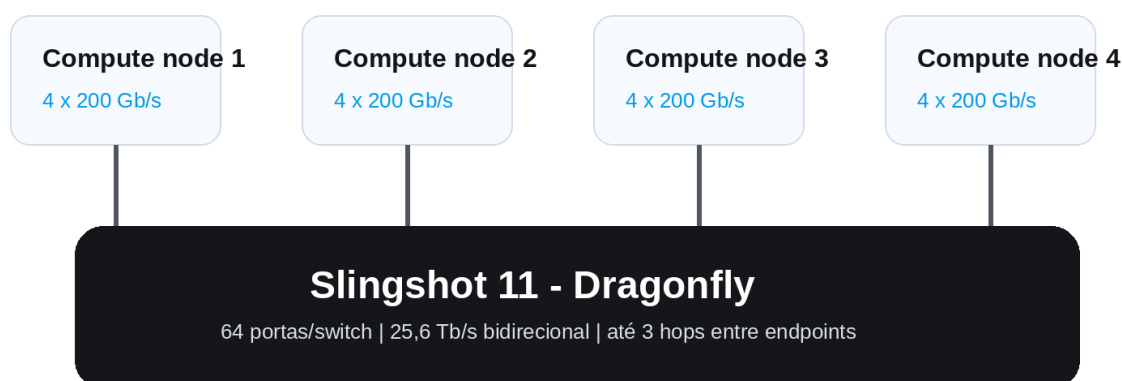
A configuração total publicada pelo LLNL inclui 32 login nodes, 11.424 batch nodes e 64 debug nodes, totalizando 11.520 nós. Cada nó de computação usa quatro MI300A. A plataforma declara 1.105.920 cores de CPU e 5.898.240 GiB de memória total.

Benchmark x sistema instalado

O TOP500 de junho de 2026 lista 11,34 milhões de "cores" para a configuração medida. Esse número segue a metodologia de contagem do benchmark e não deve ser comparado diretamente aos 1.105.920 cores x86 publicados pelo LLNL. Para engenharia de capacidade, use a definição de cada fonte e documente o escopo.

HPE Slingshot 11: o sistema nervoso do cluster

HPE Slingshot 11: fabric de baixa latência



Injeção por nó: 100 GB/s agregados. Topologia Dragonfly reduz latência e a quantidade de enlaces ópticos.

Topologia simplificada. Fonte de parâmetros: HPC @ LLNL e HPE.

O Slingshot 11 utiliza uma topologia Dragonfly. O objetivo é reduzir o número de saltos entre endpoints e, ao mesmo tempo, evitar uma explosão no uso de fibra óptica. Switches de alta radix formam grupos com conexões locais de cobre; enlaces ópticos ligam grupos distantes.

Parâmetro	Valor público
Portas por switch	64
Bandwidth bidirecional por switch	25,6 Tb/s
Porta	até 200 Gb/s unidirecional
Interfaces por compute node	4 x 200 Gb/s
Injeção agregada por nó	100 GB/s
Topologia	Dragonfly

Storage: mover dados na velocidade da computação

Hierarquia de armazenamento: Rabbit + Lustre



Fluxo típico: compute -> Rabbit -> Lustre

Hierarquia de armazenamento do El Capitan.

O El Capitan introduziu o Rabbit, uma camada de armazenamento near-node projetada para reduzir latência de I/O, absorver bursts e diminuir interferência no fabric principal. Cada módulo Rabbit combina SSDs e um processador AMD EPYC dedicado. A configuração atual do portal de hardware lista 720 módulos.

A persistência global usa Lustre. O filesystem /p/lustre4, associado ao El Capitan, é publicado atualmente com 359 PiB de capacidade - aproximadamente 400 PB em unidades decimais. Material de arquitetura do LLNL documenta throughput superior a 2,6 TB/s no tier global e uma camada Rabbit na ordem de 21 PB.

Padrão arquitetural

Para clusters de IA, o lesson learned é claro: storage não deve ser dimensionado apenas em TB ou PB. IOPS, throughput agregado, metadata, checkpoints, burst buffers e distância do dado até a GPU determinam o tempo de solução.

CAPÍTULO 08

Software stack: TOSS 4, Flux, ROCm e MPI

Hardware exascale exige uma pilha de software preparada para escala. O El Capitan usa TOSS 4 (Tri-Lab Operating System Stack), baseado em Red Hat Enterprise Linux e adaptado para integrar MI300A, Slingshot, armazenamento Rabbit e operação de larga escala.

Camada	Tecnologia / função
Sistema operacional	TOSS 4
Scheduler	Flux
GPU runtime	AMD ROCm / HIP
MPI	HPE Cray MPI
Math libraries	AMD rocBLAS e ecossistema HPC
Filesystem paralelo	Lustre
Programação	C/C++, Fortran, Python e modelos heterogêneos

O Flux é particularmente relevante porque foi criado no ecossistema LLNL para orquestrar workloads complexos e hierárquicos. Em clusters de IA corporativos, o equivalente é pensar em scheduler, containers, observabilidade, quotas e data orchestration como componentes de primeira classe - e não como itens adicionados depois.

IA + HPC: por que El Capitan é mais do que um benchmark

O El Capitan foi projetado para simulações de alta fidelidade e para novos workflows que combinam física, análise de dados e inteligência artificial. O LLNL cita aplicações em descoberta de materiais, otimização de design, manufatura avançada, digital twins e assistentes de IA treinados em dados classificados.



Instalação e manutenção de compute blades do El Capitan. Foto: LLNL.

Convergência HPC + AI

Em HPC tradicional, o objetivo é acelerar simulação numérica. Em AI, o objetivo é alimentar aceleradores com matrizes e dados em escala. No El Capitan, a mesma infraestrutura atende os dois mundos: HBM3, aceleradores, fabric de baixa latência, storage paralelo e precisão mista.

Consumo elétrico: leitura correta dos números

Consumo: três números que não devem ser misturados



Se 29,685 MW fossem mantidos 24x7 por um ano: ~260 GWh. A 36 MW constantes: ~315 GWh/ano.

Esses valores anualizados são cenários matemáticos; a carga real varia com jobs, disponibilidade e operação.

Potência de benchmark, potência de pico do sistema e capacidade do prédio representam escopos diferentes.

Na lista TOP500 de junho de 2026, o consumo medido associado ao HPL é 29.684,62 kW. Dividindo o desempenho HPL de 1,809 exaFLOP/s por essa potência, obtemos a eficiência publicada de aproximadamente 60,94 GFLOP/s por watt.

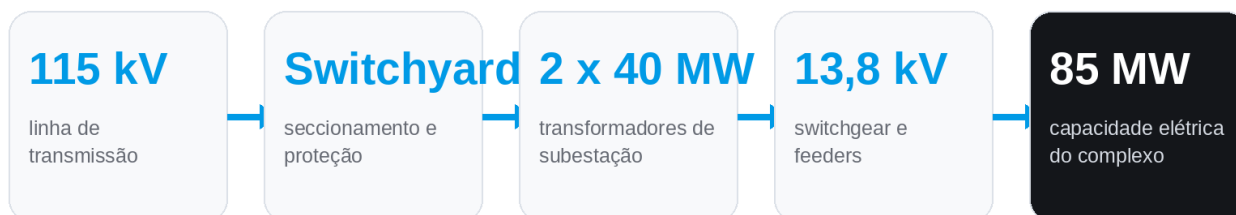
O hardware overview do LLNL publica 36,0 MW de peak power para a configuração total e 90 compute cabinets. O artigo de siting do ASC usa a ordem de grandeza de ~35 MW. Esses valores são coerentes com a densidade de cerca de 400 kW por rack divulgada pelo LLNL em 2026.

Cenário	Potência	Energia se 24x7 por 1 ano
HPL TOP500	29,685 MW	~260,0 GWh
Peak LLNL	36,0 MW	~315,4 GWh

Os valores anualizados acima são cálculos de engenharia para dimensionamento de ordem de grandeza. A operação real depende de carga, manutenção, disponibilidade, perfil dos jobs e eficiência do data center.

Do grid ao rack: a infraestrutura elétrica

Cadeia elétrica pública do ECFM



Importante: as fontes públicas não detalham toda a topologia de UPS, geração de emergencia ou segurança da instalação.

Cadeia funcional baseada nos elementos publicamente divulgados pelo ECFM.

Para receber sistemas exascale, o LLNL executou o Exascale Computing Facility Modernization (ECFM). A obra levou uma linha de transmissão de 115 kV até um novo switchyard, instalou dois transformadores de subestação de 40 MW, switchgear de 13,8 kV, relays, feeders e subestações secundárias dentro da instalação.

A capacidade elétrica do computing floor foi elevada de 45 MW para 85 MW. O objetivo do projeto não era apenas energizar o El Capitan, mas preparar o complexo para operar duas máquinas exascale simultaneamente em ciclos futuros.

O que não está sendo inferido

As fontes abertas não publicam todos os detalhes de UPS, grupos geradores, autonomia, seletividade, classificação de redundância ou controles de segurança. Este e-book não preenche essas lacunas com suposições.

CAPÍTULO 12

O data center: Building 453 e a modernização ECFM



Building 453 do Livermore Computing Complex. Fonte: LLNL.

O Building 453 foi projetado para várias gerações de supercomputadores. O LLNL informa 48.000 ft² de computer room floor - pouco mais de 4.450 m² - além de infraestrutura mecânica e elétrica de escala de utility. O prédio recebeu certificação LEED Gold em 2010.

48.000 ft²

machine floor

85 MW

capacidade elétrica ECFM

28.000 TR

capacidade de cooling

115 kV

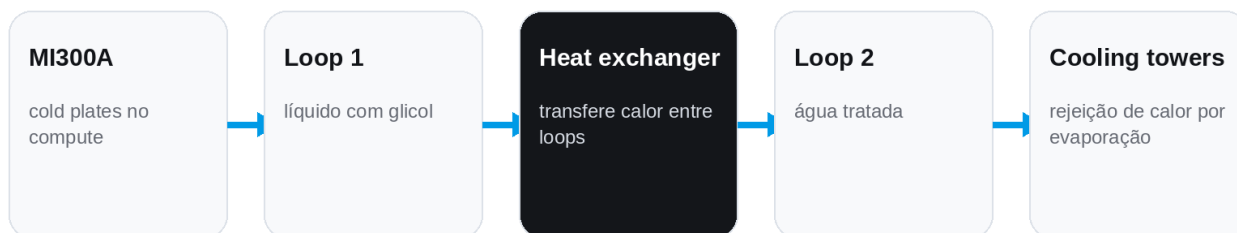
alimentação de alta tensão

Em projetos de IA, essa é a principal mudança de paradigma: o data center deixa de ser uma sala que hospeda servidores e passa a ser uma planta eletromecânica cuja carga principal é computação de alta densidade.

CAPÍTULO 13

Direct Liquid Cooling: 400 kW por rack

Liquid cooling: do chip até a atmosfera

**~400 kW/rack**

densidade citada pelo LLNL para o El Capitan

O LLNL informa que acima de ~25 kW/rack o ar deixa de ser prático para esta classe de densidade.

Fluxo térmico simplificado do El Capitan.

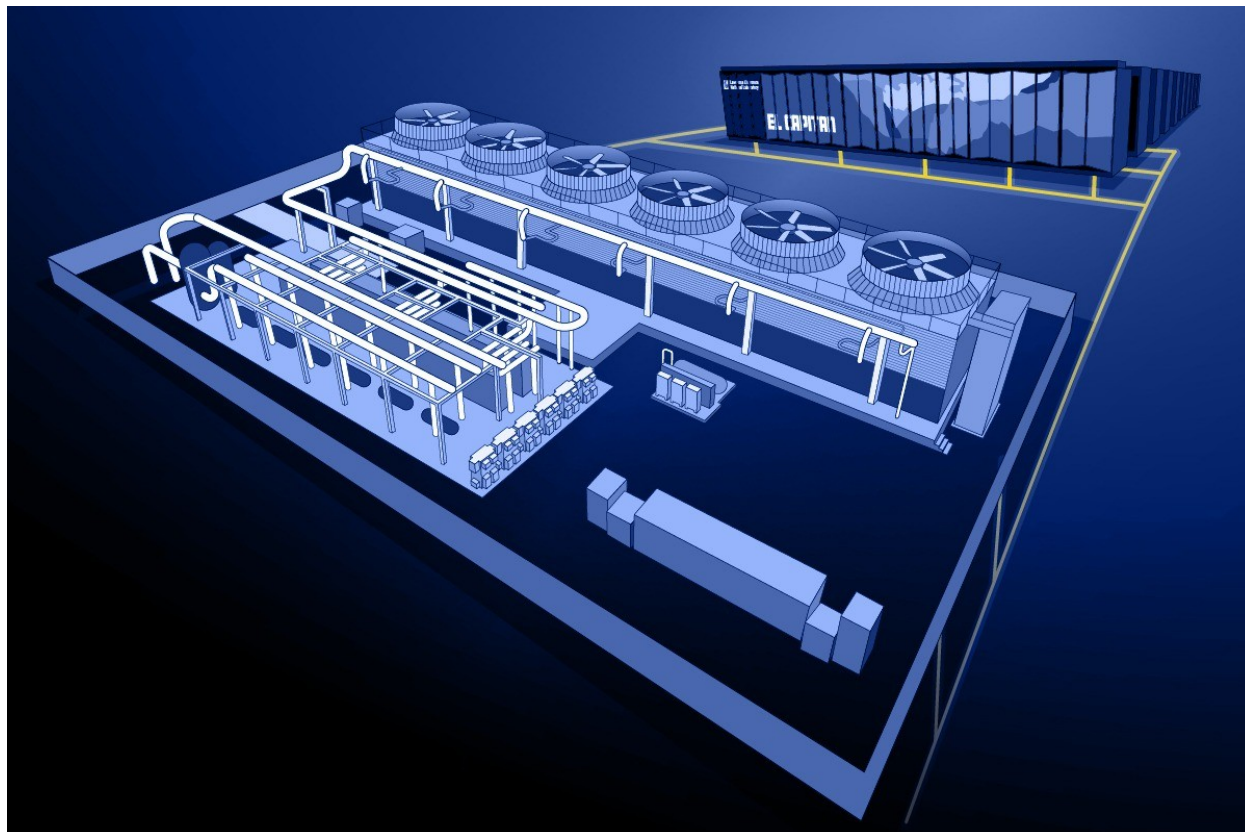
Em junho de 2026, o LLNL publicou que os racks do El Capitan atingem densidades de aproximadamente 400 kW. A equipe de facilities cita ~25 kW/rack como ponto a partir do qual o ar deixa de ser uma solução prática para essa classe de arquitetura.

O HPE Cray EX usa cold plates diretamente conectadas aos componentes de maior dissipação. O sistema é descrito pela HPE como 100% fanless e direct liquid-cooled. A eliminação de ventiladores nos compute cabinets reduz energia parasita e permite concentrar muito mais compute por metro quadrado.

Equação térmica

Quase toda a potência elétrica consumida pelo hardware termina como calor. Um sistema de 36 MW de TI pode gerar uma carga térmica da mesma ordem de grandeza. Convertendo 36 MW para toneladas de refrigeração, temos aproximadamente 10.236 TR.

Dois loops, trocadores de calor e torres



Sistema de cooling do El Capitan: torres, chillers, bombas, trocadores, sensores e tubulação. Gráfico: Dan Herchek/LLNL.

O LLNL descreve dois circuitos: um com água tratada e outro com fluido à base de glicol. Sensores acompanham temperatura, vazão e pressão, e os controles respondem a mudanças rápidas de carga computacional. A instalação utiliza mais de 2.000 ft de tubulação nessa cadeia de cooling.

A água circula em temperatura relativamente elevada, em torno de 85 °F (~29,4 °C), reduzindo a necessidade de refrigeração mecânica intensiva. A rejeição final de calor ocorre principalmente por evaporação nas torres do ECFM.

A planta térmica em escala industrial



Torres de cooling e heat exchanger do LCC B453. Fotos: LLNL.

O ECFM ampliou a capacidade de water cooling do complexo de 10.000 para 28.000 toneladas de refrigeração. Isso equivale a aproximadamente 98,5 MW de capacidade térmica nominal. O projeto adicionou seis torres de 3.000 TR, além de bombas, tubulação e trocadores.

Capacidade de cooling não é consumo de TI

Os 28.000 TR pertencem ao complexo e foram dimensionados para suportar a evolução da instalação, inclusive coexistência de máquinas exascale. Não é correto interpretar esse número como "o El Capitan consome 98,5 MW de cooling".

CAPÍTULO 16

Energia por rack e implicações para o projeto

A relação 90 compute cabinets x ~400 kW/rack produz 36 MW, exatamente a ordem de grandeza publicada como peak power da plataforma. Em um data center convencional, 400 kW em um único rack muda praticamente tudo: barramento, cabos, proteção, distribuição, peso, piping, CDU, layout e segurança operacional.

Item	Data center tradicional	Classe El Capitan / AI Factory
Densidade por rack	5-20 kW típicos	~400 kW/rack no El Capitan
Cooling	Ar / containment	Direct liquid cooling
Distribuição	PDU's convencionais	Alta corrente + arquitetura específica
Rede	10/25/100G	200G por porta, multi-rail
Storage	Capacidade predominante	Capacidade + TB/s + near-node
Operação	Carga relativamente estável	Salto de carga e calor por jobs

Para engenheiros de IA, a conclusão é direta: clusters de aceleradores devem ser projetados de forma conjunta com elétrica e mecânica. Tentar "encaixar" um cluster de centenas de kW por rack em uma sala originalmente projetada para 10 kW/rack normalmente cria gargalos de distribuição e cooling.

Contexto global: TOP500 em junho de 2026

TOP500 - junho de 2026

#	Sistema	HPL Rmax (EF/s)	Power (MW)
1	LineShine	2,1984	42,220
2	El Capitan	1,8090	29,685
3	Frontier	1,3530	24,607
4	Aurora	1,0120	38,698
5	JUPITER Booster	1,0000	--

Valores de potência conforme dados publicados no TOP500 quando disponíveis. El Capitan é o foco deste e-book.

Contexto de ranking. Fonte principal: TOP500 June 2026.

El Capitan liderou o TOP500 de novembro de 2024 até novembro de 2025. Na 67ª edição, de junho de 2026, LineShine estreou em primeiro lugar com 2,1984 exaFLOP/s. El Capitan passou para a segunda posição mantendo 1,8090 exaFLOP/s.

Para o objetivo deste e-book, a mudança de ranking não reduz o valor técnico do case. O El Capitan continua sendo uma das arquiteturas mais documentadas para estudar convergência entre HPC, IA, alta densidade elétrica, storage paralelo e liquid cooling.

O que o El Capitan ensina para data centers de IA

- Compute não pode ser dimensionado isoladamente: GPU, memória, fabric e storage formam um sistema.
- A densidade elétrica se torna o principal parâmetro de layout: 400 kW/rack muda a engenharia da sala.
- Liquid cooling deve entrar no projeto conceitual, não como retrofit de última hora.
- A energia deve ser contratada e entregue em escala de utility; 30-40 MW de TI já exigem subestação e planejamento de longo prazo.
- Near-node storage e filesystems de centenas de PB mostram que I/O é parte do desempenho computacional.
- O fabric precisa ser calculado por injeção, bisseção, latência e topologia - não apenas pela velocidade nominal da porta.
- Telemetria de energia, temperatura, vazão e pressão precisa acompanhar a dinâmica dos workloads.
- Capacidade do prédio deve incluir margem de crescimento e coexistência de gerações de hardware.

Visão EnQ Digital

Uma AI Factory é a combinação de data center, GPU, rede, storage, software e energia. O valor de negócio aparece quando essas camadas são tratadas como uma única arquitetura.

CAPÍTULO 19

Modelo de referência para um projeto corporativo de GPU

O El Capitan opera em uma escala excepcional, mas os princípios podem ser traduzidos para clusters corporativos. Um projeto de 1 a 10 MW de GPU deve iniciar pelo workload e derivar potência, cooling, rede e storage.

Etapa	Pergunta de engenharia
Workload	Treinamento, inferência, HPC, RAG, vídeo ou simulação?
GPU	Quantas GPUs, TDP, memória e interconexão intra-node?
Rack	Qual densidade máxima em kW/rack e peso por rack?
Cooling	DLC, CDU in-row/in-rack, temperatura de supply e heat rejection?
Energia	MW de TI, headroom, subestação, distribuição e expansão?
Fabric	Ethernet/RoCE ou InfiniBand; 200/400/800G; oversubscription?
Storage	Throughput, metadata, checkpoint, object e archive?
Operação	Scheduler, observabilidade, tenancy, segurança e automação?

A primeira conta deve ser feita em kW por rack e MW por sala. A segunda, em GB/s por GPU e por cluster. A terceira, em TB/s de storage. Somente depois faz sentido discutir quantidade total de racks e área.

Checklist técnico para AI Factory

- Definir envelope de potência IT por rack e por sala.
- Definir temperatura, vazão e qualidade do fluido no circuito de liquid cooling.
- Dimensionar CDUs, heat exchangers e capacidade de rejeição de calor.
- Validar disponibilidade de energia no ponto de conexão e prazo de energização.
- Modelar step loads de GPU e resposta do sistema elétrico/mecânico.
- Definir rede east-west com latência e bisection bandwidth adequadas.
- Dimensionar storage por throughput e metadata, não apenas por capacidade.
- Planejar observabilidade: power, temperatura, flow, pressure, GPU health e fabric.
- Definir crescimento por fases sem bloquear manutenção ou expansão.
- Validar compatibilidade entre servidor, rack, CDU, manifold, fluido e facility water.

Regra prática

O data center para IA deve ser pensado como uma máquina térmica e elétrica que hospeda computação, e não como uma sala de TI tradicional.

APÊNDICE

Glossário essencial

Termo	Definição
exaFLOP/s	10 ¹⁸ operações de ponto flutuante por segundo.
Rmax	Desempenho efetivamente medido no HPL.
Rpeak	Pico teórico calculado da plataforma.
HPCG	Benchmark que enfatiza memória e comunicação.
HPL-MxP	Benchmark de precisão mista, relevante para convergência HPC/AI.
APU	Pacote que integra CPU e GPU/acelerador com memória compartilhada.
HBM3	High Bandwidth Memory de terceira geração.
DLC	Direct Liquid Cooling, com liquid cooling próximo ao componente.
CDU	Cooling Distribution Unit.
Dragonfly	Topologia de interconexão para clusters de grande escala.
Lustre	Filesystem paralelo usado em HPC.
Rabbit	Near-node local storage desenvolvido para o ecossistema El Capitan.
TOSS	Tri-Lab Operating System Stack.
Flux	Resource manager e scheduler desenvolvido no LLNL.

REFERÊNCIAS

Fontes técnicas principais

TOP500 - El Capitan system page: <https://www.top500.org/system/180307/>

TOP500 - June 2026: <https://www.top500.org/lists/top500/2026/06/>

HPC @ LLNL - El Capitan: <https://hpc.llnl.gov/hardware/compute-platforms/el-capitan>

HPC @ LLNL - Hardware Overview: <https://hpc.llnl.gov/documentation/user-guides/using-el-capitan-systems/hardware-overview>

LLNL - Introducing El Capitan: <https://str.llnl.gov/str-december-2024/introducing-el-capitan>

LLNL - El Capitan verified: <https://www.llnl.gov/article/52061/lawrence-livermore-national-laboratorys-el-capitan-verified-worlds-fastest-supercomputer>

LLNL - Powering up / ECFM: <https://www.llnl.gov/article/48101/powering-llnl-prepares-exascale-massive-energy-water-upgrade>

LLNL - How El Capitan keeps its cool: <https://www.llnl.gov/article/54496/meet-machines-matter-how-el-capitan-keeps-its-cool>

ASC - Facilities: <https://asc.llnl.gov/facilities>

ASC - El Capitan: <https://asc.llnl.gov/exascale/el-capitan>

HPE - El Capitan press release: <https://www.hpe.com/us/en/newsroom/press-release/2024/11/hewlett-packard-enterprise-delivers-worlds-fastest-direct-liquid-cooled-exascale-supercomputer-el-capitan-for-lawrence-livermore-national-laboratory.html>

HPE Cray Supercomputing EX QuickSpecs: <https://www.hpe.com/us/en/collaterals/collateral.a00094635enw.html>

HPC @ LLNL - Parallel File Systems: <https://hpc.llnl.gov/hardware/file-systems/parallel-file-systems>

CRÉDITOS

Imagens e notas editoriais

Imagens reais utilizadas neste e-book são creditadas ao Lawrence Livermore National Laboratory (LLNL) e aos respectivos autores quando informados nas páginas de origem. A capa utiliza fotografia oficial do El Capitan creditada a Garry McLeod/LLNL.

- El Capitan full: Garry McLeod/LLNL.
- MI300A / blade scale: LLNL content hub.
- Instalação de compute blades: LLNL content hub.
- Building 453, cooling towers e heat exchanger: LLNL / Advanced Simulation and Computing.
- Cooling system graphic: Dan Herchek/LLNL, 2026.

Uso editorial

Antes de redistribuição comercial em larga escala, valide os termos de uso de cada imagem na fonte oficial. Os diagramas, tabelas e gráficos assinados pela EnQ Digital foram redesenhados especificamente para este material.

Dados consultados e atualizados até 22 de agosto de 2026. Rankings e especificações de supercomputadores podem mudar com novas medições, upgrades e novas edições do TOP500.



Infraestrutura para a era da Inteligência Artificial.

Cloud - IA - Data Center - Conectividade

EnQ Digital

enqdigital.com.br