

Da História às Fábricas de IA

Chatbots, imagem e vídeo generativos, transformação empresarial, GPUs, data centers, energia, liquid cooling e conectividade global.

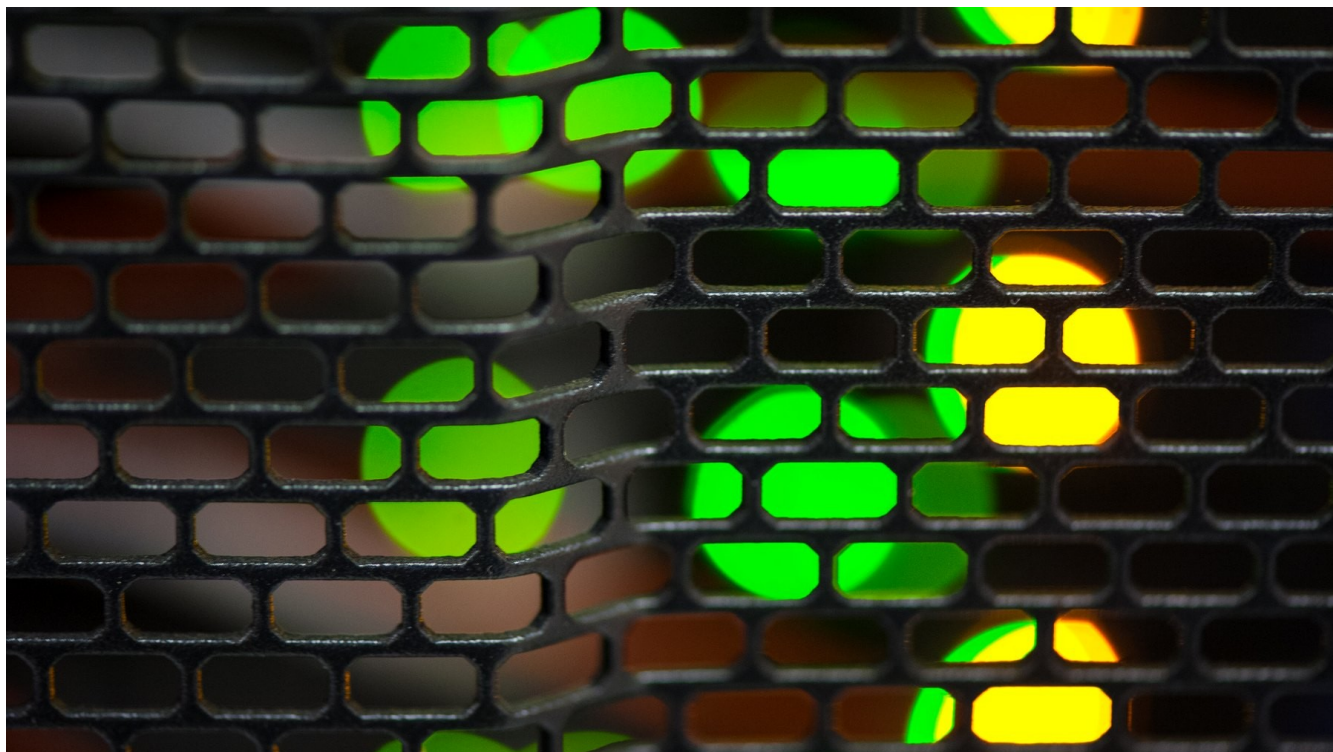


Foto real: Dennis Schroeder / NREL / U.S. Department of Energy - domínio público.

EDIÇÃO 2026

Um guia didático para entender a inteligência artificial como tecnologia, plataforma de negócios e infraestrutura física.

A IA deixou de ser apenas software

Por trás de uma resposta em segundos existe uma cadeia física de computação, energia, cooling, storage e fibra óptica.

A inteligência artificial evoluiu de experimentos acadêmicos e sistemas de regras para modelos generativos capazes de conversar, criar imagens, gerar vídeo, programar e operar processos. A partir desta virada, a discussão sobre IA deixou de caber somente no software: capacidade de GPU, energia disponível, refrigeração e conectividade passaram a determinar quanto, onde e com que custo a IA pode escalar.

US\$ 4,8 tri

mercado global de IA projetado
para 2033

945 TWh

consumo anual de data centers
projetado para 2030

1.360

data centers hyperscale no fim de
2025

86,8%

participação renovável na matriz
elétrica brasileira em 2025

>99%

dos fluxos internacionais de dados
passam por cabos submarinos

20 PFLOP/s

ordem de capacidade instalada do
Santos Dumont

A EQUAÇÃO DA IA MODERNA

Modelos + dados + computação + energia + refrigeração + conectividade + capital. Se um desses componentes vira gargalo, a escala da IA também vira gargalo.

Fonte: UNCTAD (2025); IEA (2025); Synergy Research Group (2026); EPE (2026); ITU (2025/2026); LNCC (consulta em 2026).

Do algoritmo ao megawatt

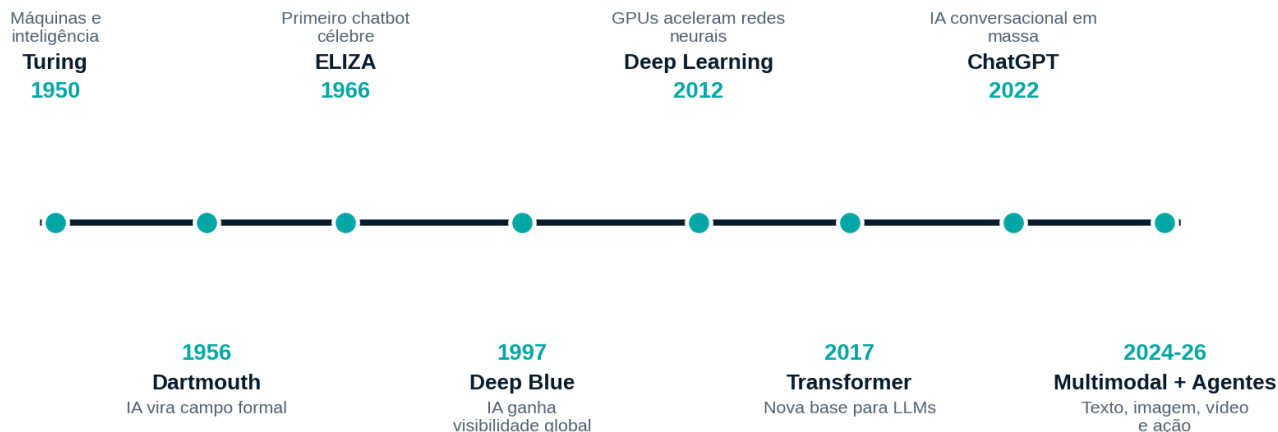
- 01** **História e evolução**
Das ideias de Turing ao deep learning
- 02** **Chatbots, LLMs e agentes**
De ELIZA a sistemas que usam ferramentas e executam tarefas
- 03** **Imagem, vídeo e multimodalidade**
Quando a IA passa a produzir e interpretar várias mídias
- 04** **Empresas e economia**
Produtividade, automação, novos modelos de negócio e capital
- 05** **Infraestrutura de IA**
GPU, memória, rede, storage e AI Factories
- 06** **Liquid cooling**
Por que racks de alta densidade mudam a engenharia térmica
- 07** **Case Santos Dumont**
H100, GH200, MI300A, InfiniBand e DLC no Brasil
- 08** **Energia para IA**
Rede, PPAs, solar, eólica, hidráulica, BESS e resiliência
- 09** **Status global de data centers**
Hyperscale, colocation, enterprise e crescimento acelerado
- 10** **Conectividade global**
DCI, DWDM, fibra e cabos submarinos
- 11** **Brasil na nova geografia digital**
Energia renovável, conectividade e oportunidade de infraestrutura
- 12** **O que vem depois**
AI Factories, soberania computacional e agentes autônomos

A inteligência artificial nasceu antes da infraestrutura para executá-la

As ideias amadureceram por décadas; dados, redes neurais e GPUs destravaram a escala.

Em 1950, Alan Turing ajudou a transformar a pergunta "máquinas podem pensar?" em um problema científico. Em 1956, o encontro de Dartmouth consolidou o termo Artificial Intelligence e organizou um campo de pesquisa dedicado à linguagem, raciocínio, abstração e aprendizado por máquinas.

Da pergunta científica a uma infraestrutura econômica global



Linha do tempo resumida da IA.

Infográfico EnQ Digital. Marcos históricos sintetizados a partir de literatura acadêmica e institucional.

POR QUE HOUVE PERÍODOS DE FRUSTRAÇÃO?

Muitas ideias eram boas, mas o hardware, os dados e o custo computacional ainda eram insuficientes. Os chamados "AI winters" lembram que inovação de algoritmo sem infraestrutura pode não chegar à escala.

Da IA simbólica ao deep learning

1. Regras e sistemas especialistas

A primeira grande escola de IA tentou representar conhecimento em regras explícitas: se uma condição ocorre, então uma ação deve ser executada. Funcionava bem em domínios controlados, mas o mundo real tem exceções, ambiguidades e combinações demais.

2. Machine Learning

Com mais dados e capacidade de processamento, os sistemas passaram a aprender padrões a partir de exemplos. Em vez de programar todas as regras, o desenvolvedor define dados, objetivo, método de treinamento e métricas.

3. Deep Learning + GPU

Redes neurais profundas ampliaram a capacidade de reconhecimento de imagens, voz e linguagem. GPUs, inicialmente populares em computação gráfica, tornaram-se essenciais porque executam grande volume de operações matemáticas em paralelo.

A pilha física da inteligência artificial

Aplicações e agentes	LLMs, RAG, copilotos, automação
Plataforma de IA	Frameworks, scheduler, Kubernetes, observabilidade
Compute acelerado	GPU / aceleradores / CPU / HBM
Rede + Storage	InfiniBand / Ethernet HPC / NVMe / Lustre
Energia + Cooling	Subestação, UPS, DLC, CDU, rejeição de calor
Data center + conectividade	Racks, segurança, DCI, fibra e cabos submarinos

A IA moderna é uma pilha que termina no data center.
Infográfico EnQ Digital.

De ELIZA aos assistentes generativos

A interface conversacional evoluiu de regras para modelos que trabalham com contexto, linguagem, código e ferramentas.

ELIZA, desenvolvido por Joseph Weizenbaum no MIT nos anos 1960, é um marco da história dos chatbots. Seu script DOCTOR respondia por padrões e transformações de texto. O sistema não "entendia" como um modelo moderno, mas demonstrou o quanto uma conversa textual pode induzir percepção de inteligência.

```
Welcome to

EEEEEE LL      IIII  ZZZZZZ  AAAAA
EE      LL      II   ZZ     AA   AA
EEEEEE LL      II   ZZ     AAAAAA
EE      LL      II   ZZ     AA   AA
EEEEEE LLLLLL IIII ZZZZZZ  AA   AA

Eliza is a mock Rogerian psychotherapist.
The original program was described by Joseph Weizenbaum in 1966.
This implementation by Norbert Landsteiner 2005.

ELIZA: Is something troubling you ?
YOU:   Men are all alike.
ELIZA: What is the connection, do you suppose ?
YOU:   They're always bugging us about something or other.
ELIZA: Can you think of a specific example ?
YOU:   Well, my boyfriend made me come here.
ELIZA: Is it important to you that your boyfriend made you come here ?
YOU:   He says I'm depressed much of the time.
ELIZA: I am sorry to hear that you are depressed.
YOU:   It's true. I am unhappy.
ELIZA: Can you explain what made you unhappy ?
YOU:
```

Exemplo de conversa com uma implementação do ELIZA.

Imagem real/reprodução de terminal: Wikimedia Commons - texto em domínio público.

O QUE MUDOU?

Bots antigos seguiam fluxos. LLMs calculam continuações prováveis com base em representações aprendidas em larga escala e podem combinar contexto, instruções e ferramentas externas.

LLMs, copilotos e agentes

A arquitetura Transformer, publicada em 2017, criou um caminho eficiente para modelar dependências em sequências. Grandes modelos de linguagem ampliaram esse princípio com escala de dados e computação. A popularização de interfaces como o ChatGPT, a partir de 2022, fez a IA conversacional chegar ao grande público.

A evolução da interface: responder -> raciocinar -> executar



O salto de produtividade acontece quando a IA deixa de ser apenas uma interface e passa a operar processos com governança.

Da resposta automática à execução de processos.

Infográfico EnQ Digital.

O que caracteriza um agente de IA?

- Recebe um objetivo e decompõe a tarefa em etapas.
- Consulta sistemas, APIs, bancos de dados ou bases RAG autorizadas.
- Usa ferramentas, executa ações e registra resultados.
- Pode trabalhar com outros agentes especializados.
- Exige identidade, permissões, auditoria, limites e aprovação humana quando o risco aumenta.

PONTO DE GOVERNANÇA

Quanto mais autonomia, maior deve ser o controle de acesso, rastreabilidade, avaliação de risco e separação entre recomendação e execução.

Texto, imagem, voz e vídeo convergem

A IA deixou de apenas classificar mídias e passou a criá-las.

Na visão computacional tradicional, a pergunta era: "o que existe nesta imagem?". Com modelos generativos, a pergunta virou: "que imagem ou cena deve ser criada a partir desta descrição?". Modelos de difusão e arquiteturas multimodais reduziram a distância entre linguagem natural e produção visual.

Modalidade	Antes	Agora
Texto	Classificação e busca	Redação, síntese, código, análise e diálogo
Imagem	Reconhecimento de objetos	Geração, edição, variações e composição
Áudio	Transcrição e comandos	Voz sintética, clonagem autorizada, tradução e diálogo
Vídeo	Análise de frames	Geração de cenas, edição e pré-visualização
Multimodal	Pipelines separados	Texto, imagem, áudio e vídeo no mesmo contexto

IMPACTO CRIATIVO

Marketing, design, games, audiovisual, educação e e-commerce passaram a usar linguagem natural como interface de criação. O diferencial competitivo muda de "acessar a ferramenta" para construir processos, dados, identidade e curadoria melhores.

Por que vídeo custa mais computação?

Vídeo acrescenta uma dimensão temporal. O modelo precisa manter coerência entre frames, movimento, câmera, iluminação e identidade visual. Isso eleva memória, throughput e tempo de inferência - e pressiona ainda mais a infraestrutura de GPU.

A transformação acontece quando IA entra no processo

O ganho real aparece quando modelo, dados e operação são integrados com governança.

Área	Aplicações de IA	Valor esperado
Comercial	Qualificação, propostas, pesquisa, follow-up	Mais velocidade e cobertura
Marketing	Conteúdo, segmentação, imagem, vídeo, análise	Mais experimentação e personalização
Atendimento	Assistente, busca em conhecimento, resumo	Menor tempo de resposta
Tecnologia	Código, testes, documentação, troubleshooting	Maior produtividade técnica
Operações	Previsão, manutenção, visão computacional	Menos falhas e melhor utilização
Financeiro	Conciliação, documentos, anomalias, forecast	Mais controle e previsibilidade
Gestão	BI conversacional, cenários, copilotos executivos	Decisões mais informadas

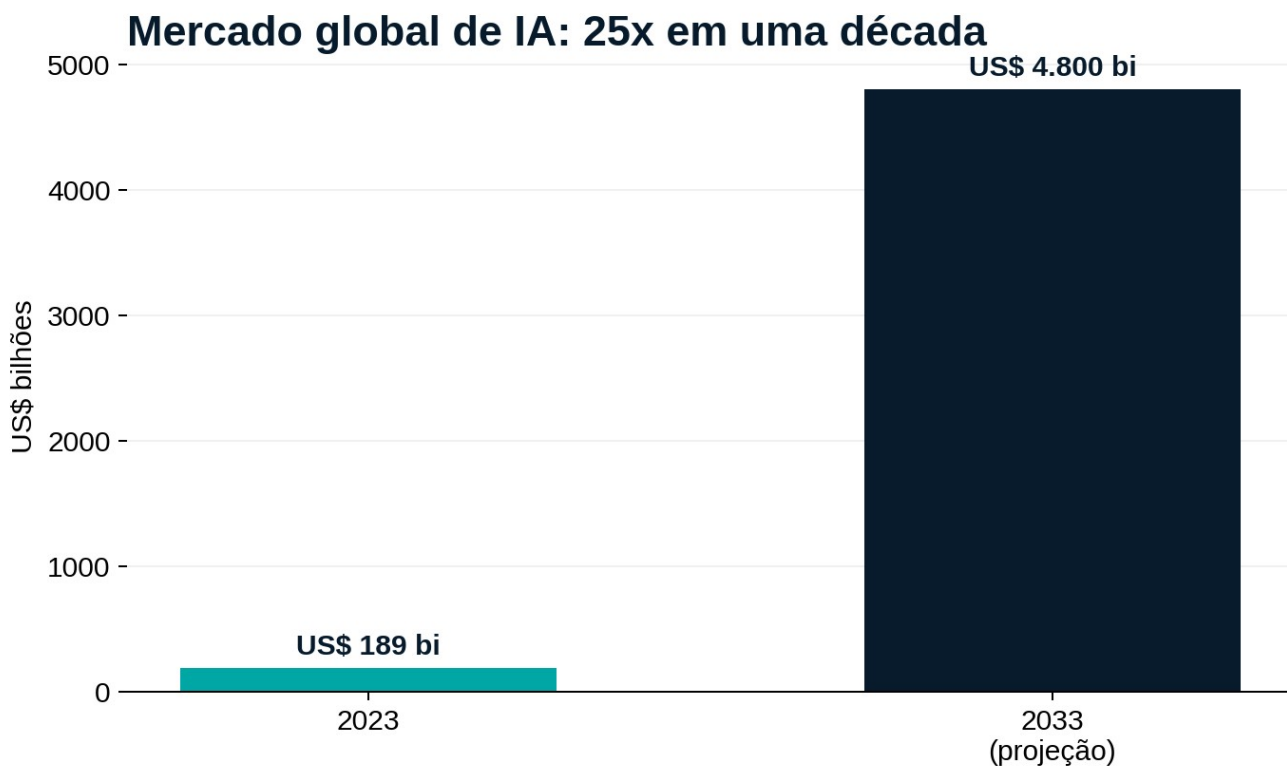
A REGRA PRÁTICA

Comece pelo processo, não pelo modelo. Defina KPI, dados, risco, integração e responsável. Depois escolha a tecnologia.

A IA corporativa tende a migrar de ferramentas individuais para uma camada operacional: copilotos em sistemas existentes, agentes conectados a CRM/ERP/ITSM, automação de documentos, análise preditiva e interfaces conversacionais sobre dados proprietários.

Uma nova corrida por computação e capital

O crescimento da IA movimenta uma cadeia física inteira - de semicondutores a usinas e cabos.



Fonte: UNCTAD, Technology and Innovation Report 2025. 2033 é uma projeção.

Projeção de mercado global de IA segundo a UNCTAD.

Gráfico EnQ Digital com dados da UNCTAD, Technology and Innovation Report 2025.

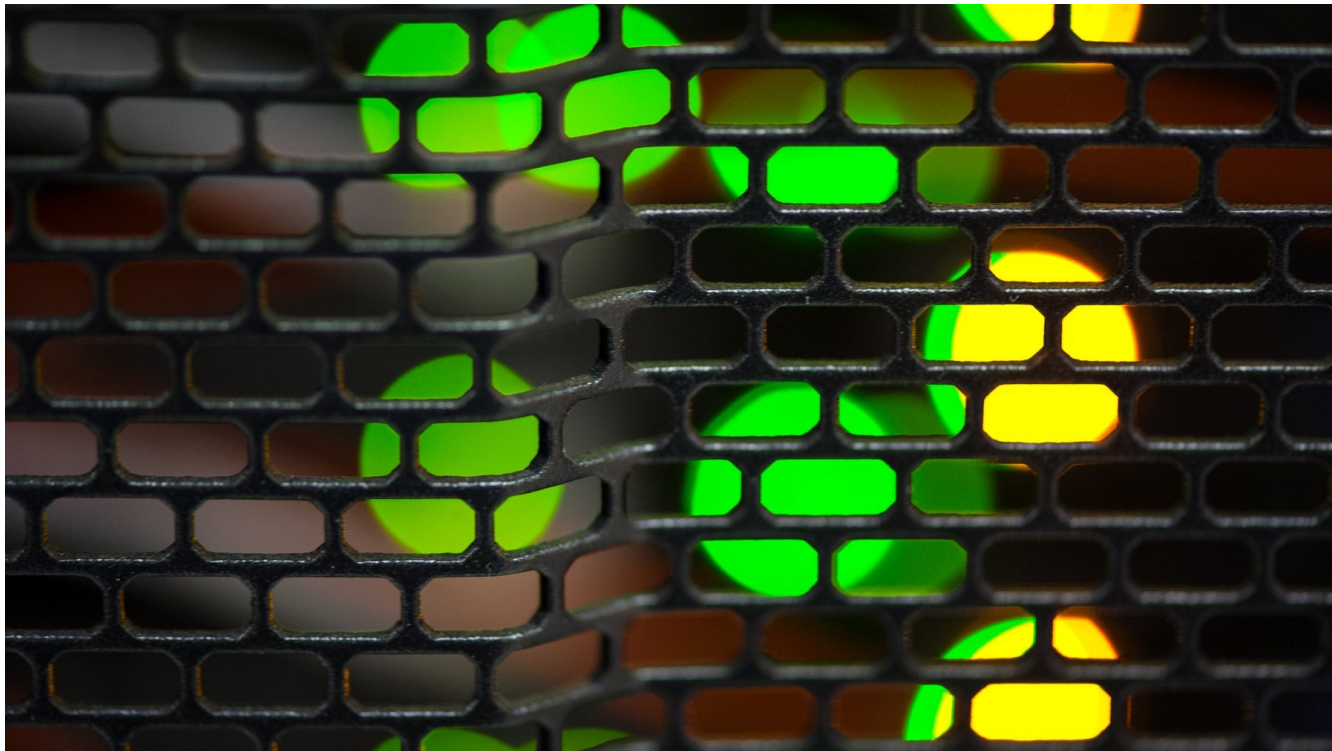
A UNCTAD projeta que o mercado global de IA avance de US\$ 189 bilhões em 2023 para US\$ 4,8 trilhões em 2033. Trata-se de uma projeção, não de garantia, mas a ordem de grandeza mostra por que chips, data centers, energia e fibra passaram a receber investimentos extraordinários.

A CADEIA DA NOVA ECONOMIA

Semicondutores -> GPU e HBM -> servidores -> networking -> storage -> data center -> energia -> modelos -> aplicações. O valor econômico se distribui por várias camadas, e o gargalo pode migrar de uma para outra.

A nuvem tem endereço físico

Toda IA em produção termina em silício, energia, cooling e fibra.



Infraestrutura real de computação de alto desempenho no NREL, Estados Unidos.

Dennis Schroeder / NREL / U.S. Department of Energy - domínio público.

Quando um usuário gera uma resposta, uma imagem ou um vídeo, a solicitação percorre redes, entra em um data center, consome memória e compute acelerado e retorna pela infraestrutura de conectividade. Em escala de milhões de usuários, pequenas ineficiências viram grandes custos.

PENSAMENTO DE ENGENHARIA

O servidor é apenas um componente. Para IA em escala, o data center inteiro participa do desempenho: potência por rack, cooling, topologia de rede, storage, fibra, automação e operação 24x7.

Por que GPUs mudaram a IA?

Característica	CPU	GPU / acelerador
Objetivo	Versatilidade e baixa latência por tarefa	Paralelismo massivo
Núcleos	Poucos núcleos complexos	Muitos núcleos/unidades paralelas
Força	Lógica geral, sistemas, bancos de dados	Matrizes, tensores, treinamento e inferência
Memória	RAM convencional	HBM e alta largura de banda
Escala IA	Coordena e prepara workloads	Executa grande parte do compute neural

Treinamento x inferência

Treinamento ajusta os parâmetros de um modelo usando enormes volumes de dados e pode ocupar milhares de aceleradores por longos períodos. Inferência é o uso do modelo já treinado. Em serviços populares, a inferência pode se tornar o maior consumo agregado porque ocorre continuamente.

O GARGALO INVISÍVEL

Uma GPU rápida pode ficar ociosa esperando dados, outra GPU ou storage. Por isso redes de baixa latência e storage paralelo são parte da performance da IA.

O que diferencia uma AI Factory?

Alta densidade exige que compute, rede, storage, energia e cooling sejam projetados como um único sistema.

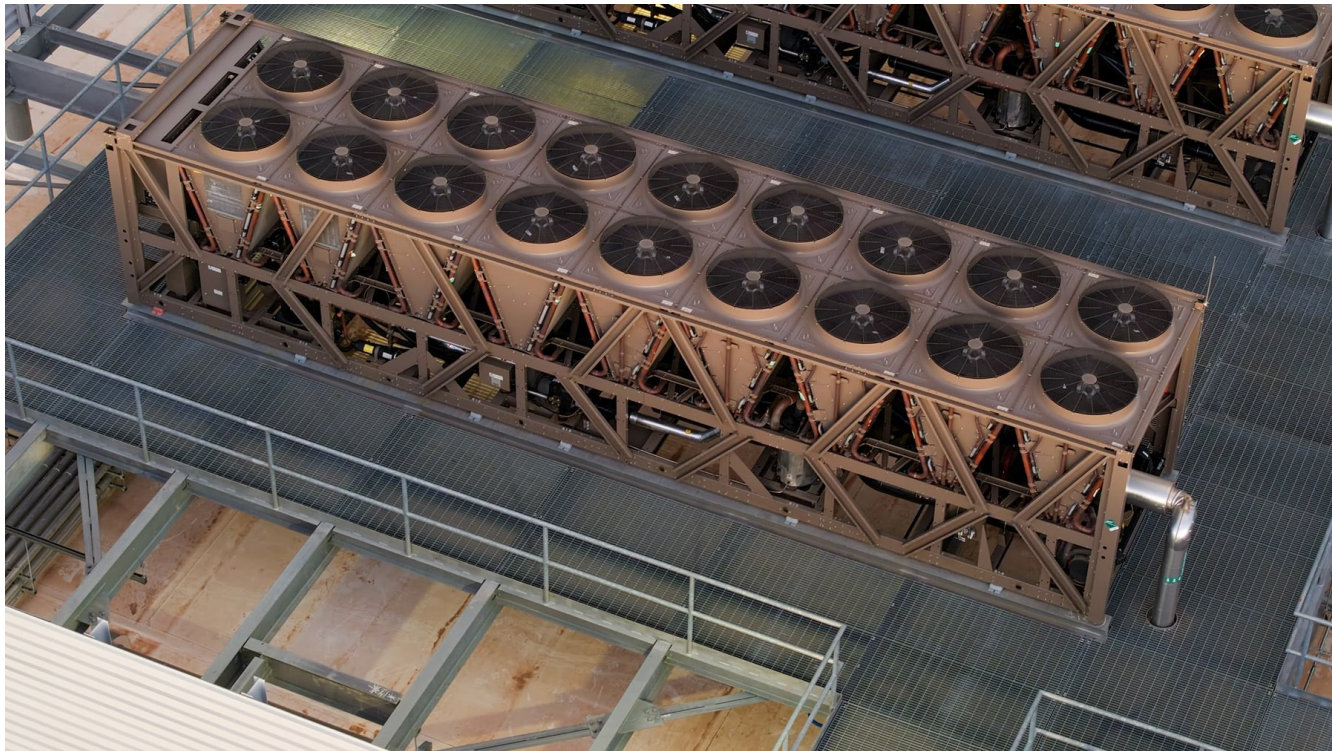
Camada	Requisito	Por que importa
Compute	GPU, CPU e aceleradores	Executar treinamento e inferência
Memória	HBM + RAM de alta capacidade	Alimentar os aceleradores sem espera
Fabric	InfiniBand ou Ethernet HPC/RDMA	Comunicar milhares de aceleradores
Storage	NVMe, paralelo, alto throughput	Datasets, checkpoints e modelos
Energia	Subestação, UPS e distribuição densa	Suportar racks de alta potência
Cooling	DLC/CDU/chillers/dry coolers	Remover calor com eficiência
DCI	Fibra, DWDM, rotas redundantes	Conectar regiões e clusters
Software	Scheduler, drivers, observabilidade	Usar o hardware de forma eficiente

AI FACTORY

O termo descreve uma infraestrutura desenhada para transformar energia + dados em capacidade de treinamento e inferência em escala. A densidade por rack pode ser várias vezes maior que em ambientes corporativos tradicionais.

O problema térmico da IA

Quanto mais compute por rack, mais calor precisa sair do mesmo volume.



Sistema real de rejeição de calor em um data center em Mesa, Arizona.

Foto: Rsparks3 / Wikimedia Commons - CC0 1.0.

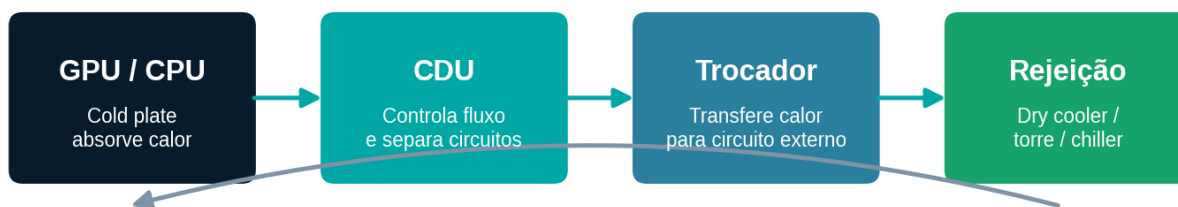
Quase toda energia elétrica consumida pelo equipamento de TI termina como calor. Em racks de alta densidade, mover esse calor apenas com ar exige grandes volumes de fluxo e ventilação. A refrigeração líquida aproxima um meio térmico muito mais eficiente da fonte de calor.

NÃO CONFUNDA

Liquid cooling não significa necessariamente consumo contínuo de grandes volumes de água. O circuito interno pode ser fechado. O consumo hídrico depende de como o calor é rejeitado externamente - por exemplo, dry cooler, chiller ou torre evaporativa.

Como funciona o Direct Liquid Cooling

Direct Liquid Cooling: o calor sai do chip pelo líquido



Circuito fechado: o fluido é recirculado; o consumo de água depende da arquitetura de rejeição de calor.

Case Santos Dumont: água de entrada 26-30 °C, saída 36-39 °C e captura de mais de 98,5% do calor (Eviden/Bull).

Fluxo simplificado de um sistema Direct Liquid Cooling.
Infográfico EnQ Digital. Parâmetros do case Santos Dumont: Eviden/Bull.

Por que o líquido é eficiente?

Líquidos transportam muito mais calor por unidade de volume do que o ar. Isso permite reduzir a dependência de ventiladores, aumentar densidade e recuperar calor em temperaturas mais úteis. A engenharia, porém, fica mais sofisticada: qualidade do fluido, pressão, vazão, redundância, detecção de vazamento, CDU e manutenção precisam entrar no projeto.

MÉTRICA MAIS IMPORTANTE

Não avalie cooling apenas pelo PUE. Para IA, observe também densidade por rack, temperatura de água, capacidade de rejeição, redundância, WUE quando aplicável e eficiência da cadeia elétrica.

Santos Dumont: IA, HPC e liquid cooling no Brasil

O supercomputador do LNCC mostra, na prática, a convergência entre compute acelerado, fabric de baixa latência, storage paralelo e engenharia térmica.



Foto: Neila Rocha/MCTIC - CC BY 2.0.

Santos Dumont (SDumont)

Localizado no LNCC, em Petrópolis-RJ, o sistema atua como Tier-0 do SINAPAD e possui capacidade instalada na ordem de 20 PFLOP/s.

A configuração atual combina plataforma BullSequana X1000 e XH3000, com diferentes gerações de CPU e aceleradores.

O projeto é relevante porque mostra como IA e supercomputação exigem arquitetura integrada: aceleradores, fabric InfiniBand, storage paralelo e cooling de alta eficiência.

62 nós

4x NVIDIA H100 SXM por nó

36 nós

4x NVIDIA GH200 por nó

18 nós

2x AMD MI300A por nó

Fonte: LNCC - Configuração do SDumont, consulta em 2026.

A arquitetura atual do Santos Dumont

Componente	Configuração publicada
BullSequana XH3420	60 nós CPU, 2x AMD Genoa-X 9684X e 1,5 TB RAM por nó
BullSequana XH3145-H	62 nós; 4x NVIDIA H100 SXM 80 GB por nó
BullSequana XH3515-H	36 nós; 4x NVIDIA GH200 Grace Hopper por nó
BullSequana G383-R80	18 nós; 2x AMD Instinct MI300A por nó
Interconexão	InfiniBand de alto rendimento e baixa latência
Storage	Lustre + DDN Exascaler; ordem de 4 PB agregados

DLC NO CASE

A Eviden/Bull informa que a expansão XH3000 usa Direct Liquid Cooling com água de entrada entre 26 °C e 30 °C e retorno entre 36 °C e 39 °C, capturando mais de 98,5% do calor de fontes, processadores, aceleradores, rede, discos e memória.

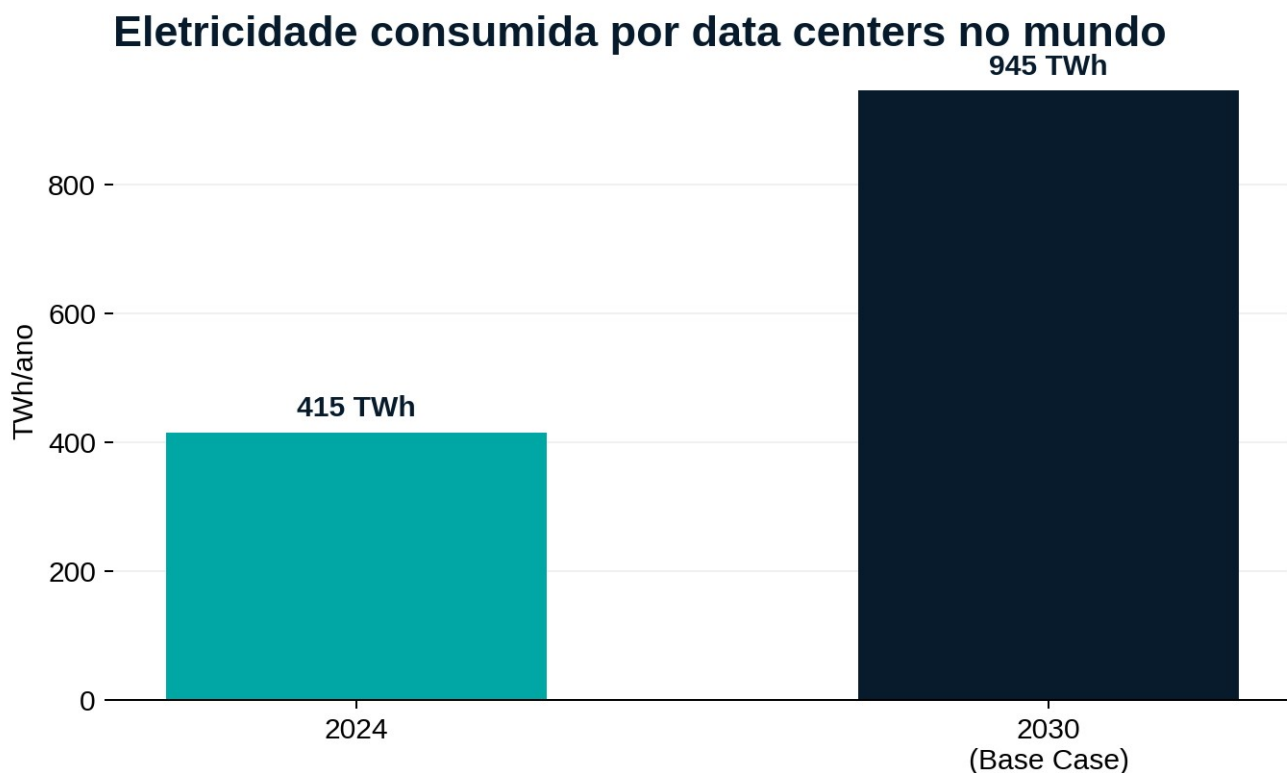
DENSIDADE

Segundo a Eviden/Bull, o sistema DLC permite densidade de até 5x a de sistemas tradicionais resfriados a ar. O case materializa a tendência central dos data centers de IA: mais compute por metro quadrado exige uma nova engenharia térmica.

Fonte: LNCC (sdumont.lncc.br/machine.php) e Eviden/Bull, press release da expansão do Santos Dumont.

O crescimento da IA vira demanda elétrica

Energia passa a ser fator de localização, prazo e custo de um cluster.



Fonte: IEA, Energy and AI (2025). 2030 = cenário central/Base Case.

Projeção da IEA para o consumo de eletricidade de data centers.

Gráfico EnQ Digital com dados da IEA, Energy and AI (2025).

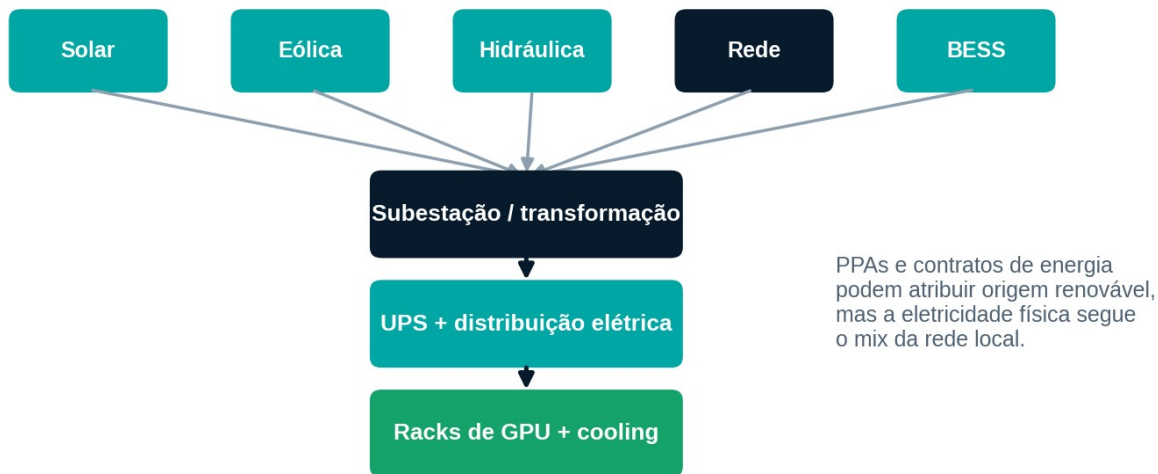
A IEA estima cerca de 415 TWh de eletricidade consumidos por data centers em 2024, aproximadamente 1,5% do consumo elétrico mundial. No cenário central, a demanda sobe para cerca de 945 TWh em 2030. A IA é apontada como o principal vetor incremental.

O GARGALO PODE ESTAR FORA DO DATA CENTER

GPUs podem ser entregues em meses; linhas de transmissão, subestações e conexões de grande porte podem levar anos. Por isso, disponibilidade de potência e prazo de energização entram cedo no estudo de viabilidade.

De onde vem a energia de um data center?

Da fonte de energia ao rack de GPU



Fluxo simplificado da energia até o rack.

Infográfico EnQ Digital.

Um data center normalmente recebe eletricidade de uma rede que combina várias fontes. A empresa pode contratar PPAs, certificados ou contratos de longo prazo, mas o elétron físico que chega ao site segue o mix do sistema elétrico local. Por isso, "origem contratual" e "mix físico" são conceitos diferentes.

Fontes e mecanismos que aparecem nos novos projetos

- Solar e eólica: construção relativamente rápida e custos competitivos em muitas regiões.
- Hidráulica: fonte firme ou flexível conforme o sistema e a hidrologia.
- BESS: baterias ajudam em flexibilidade, picos e integração de renováveis; não substituem sozinhas geração de longa duração.
- Gás, nuclear e geotermia: fontes despacháveis ou firmes relevantes em determinados mercados.
- Rede + PPAs: o modelo mais comum combina conexão robusta ao grid com contratos de energia e redundância local.

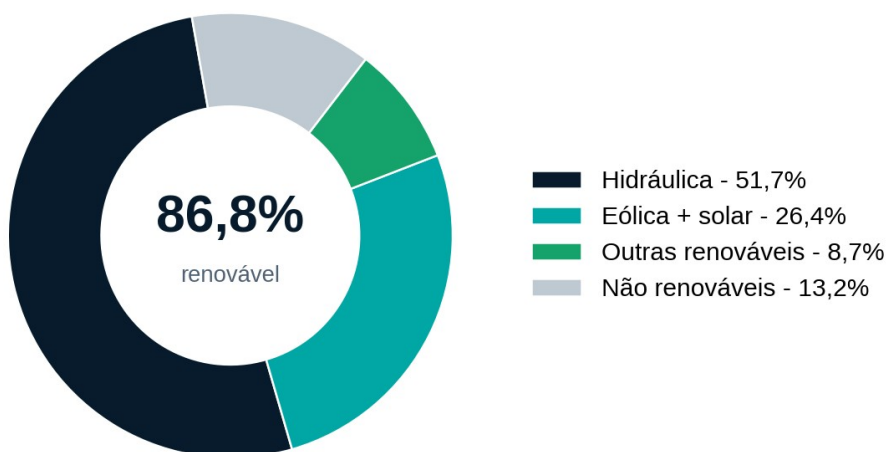
Brasil: uma matriz elétrica com forte participação renovável

A vantagem existe, mas deve ser convertida em potência disponível, transmissão, conexão e contratos executáveis.



Fotos reais: U.S. Department of Energy - domínio público (exemplos de tecnologias solar e eólica).

Matriz elétrica brasileira - 2025



Fonte: EPE, BEN/Anuário de Energia Elétrica 2026 (ano-base 2025).

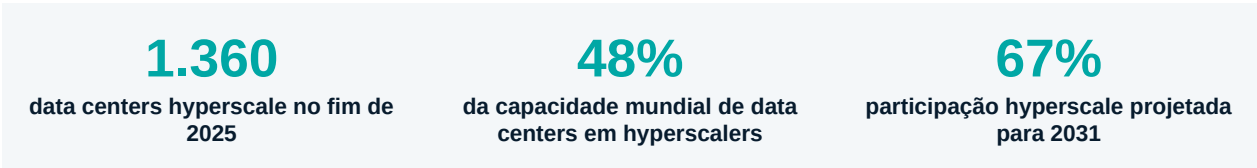
Participação das fontes na matriz elétrica brasileira em 2025.

Gráfico EnQ Digital com dados da EPE, ano-base 2025.

Segundo a EPE, as fontes renováveis responderam por 86,8% da oferta interna de energia elétrica em 2025. Hidráulica respondeu por 51,7%; eólica e solar, juntas, por 26,4%. Essa característica pode contribuir para a atratividade do Brasil em infraestrutura digital intensiva em energia.

Hyperscale se torna a camada dominante

A IA acelera uma tendência que já vinha sendo puxada por cloud e serviços digitais.



Segundo a Synergy Research Group, o número de grandes data centers operados por empresas hyperscale chegou a 1.360 no fim de 2025. Eles representavam 48% da capacidade mundial. A projeção da empresa é de que a participação chegue a 67% em 2031, enquanto a capacidade total do mercado continua crescendo.

Modelo	Papel típico	Tendência com IA
Enterprise on-premise	Sistemas internos e legado	Pode receber GPU local, mas perde participação agregada
Colocation	Infraestrutura compartilhada profissional	Cresce em MW e hospeda ambientes privados/híbridos
Hyperscale	Cloud e plataformas globais	Principal vetor de expansão de capacidade
Neocloud / GPU cloud	Compute acelerado sob demanda	Cresce com workloads de treinamento e inferência
AI Factory dedicada	Clusters massivos e alta densidade	Projeto orientado a energia, fabric e liquid cooling

Fonte: Synergy Research Group, abril de 2026.

Compute sem rede vira compute ocioso

Dentro do cluster e entre data centers, largura de banda e latência são parte do desempenho da IA.



Instalação real de fibra óptica em infraestrutura crítica.

Foto: MTA Capital Construction Mega Projects - CC BY 2.0.

Treinamentos distribuídos trocam grandes volumes de dados entre aceleradores. Sistemas de storage também precisam alimentar GPUs continuamente. Fora do cluster, DCI conecta data centers metropolitanos e regionais, enquanto redes de longa distância e cabos submarinos conectam continentes.

CAPACIDADE NÃO É A ÚNICA MÉTRICA

Para IA distribuída, observe latência, jitter, perda, oversubscription, caminho físico, redundância, MTU, RDMA, telemetria e comportamento durante falhas.

DCI, DWDM e redes de altíssima capacidade

Como data centers e clusters de IA se conectam globalmente



DCI + DWDM: múltiplos comprimentos de onda aumentam a capacidade de um par de fibras.

Baixa latência, diversidade de rotas, peering e redundância são requisitos de arquitetura - não apenas velocidade nominal.

Da empresa ao AI cluster global.

Infográfico EnQ Digital.

DCI - Data Center Interconnect

DCI conecta instalações diferentes com enlaces ópticos dedicados ou serviços de alta capacidade. Em arquiteturas distribuídas, duas unidades podem operar como partes de um mesmo ambiente lógico, respeitando limites de distância e latência.

DWDM - várias "faixas" na mesma fibra

Dense Wavelength Division Multiplexing transporta vários comprimentos de onda em um par de fibras. Cada lambda pode carregar centenas de gigabits por segundo, permitindo escalar a capacidade sem instalar um novo cabo para cada crescimento.

400G, 800G E ALÉM

A indústria óptica evolui rapidamente. A taxa nominal precisa ser analisada junto com distância, modulação, número de lambdas, budget óptico, FEC e arquitetura da rede.

Cabos submarinos: o backbone invisível do planeta



Sinalização real de infraestrutura de cabo submarino na Austrália.

Foto: ~ ItsJegh ~ / Wikimedia Commons - CC BY 2.0.

>99%

dos fluxos internacionais de dados
passam por cabos submarinos

500+

sistemas operacionais em 2025

~1,4 mil km

de extensão da rede operacional
segundo o GCR 2025

A ITU descreve os cabos submarinos como a espinha dorsal oculta da conectividade global. Eles transportam cloud, transações financeiras, conteúdo e tráfego entre regiões. A resiliência depende de diversidade de rotas, landing stations, capacidade de reparo e planejamento contra falhas físicas.

IA AUMENTA A IMPORTÂNCIA DA REDE

Modelos, datasets e usuários estão distribuídos. A capacidade computacional pode estar em outro país; sem conectividade internacional robusta, a experiência e a resiliência degradam.

Fonte: ITU Global Connectivity Report 2025 e Submarine Cable Resilience, 2026.

Energia + conectividade + demanda: a geografia da IA

O melhor local para um cluster é aquele em que vários requisitos críticos coincidem.

Pergunta de projeto	Por que decide o local
Há potência disponível?	Sem MW contratáveis, GPU não entra em produção.
Qual o prazo de conexão?	Subestação e transmissão podem definir o cronograma.
Qual a origem da energia?	Custo, emissão e compromissos ESG dependem do mix/contrato.
Há rotas ópticas diversas?	Reduz risco de falhas e melhora latência.
Qual a proximidade de cabos/IXPs/clouds?	Afeta conectividade nacional e internacional.
O site suporta liquid cooling?	Densidade e expansão dependem da infraestrutura térmica.
Há cadeia de manutenção e peças?	MTTR e disponibilidade dependem da operação local.

O Brasil combina mercado consumidor relevante, ampla infraestrutura de telecomunicações e uma matriz elétrica com forte componente renovável. São Paulo concentra grande parte da demanda e da interconexão corporativa; Fortaleza ocupa posição estratégica em rotas submarinas do Atlântico. A oportunidade, contudo, exige projetos com capacidade elétrica real, conectividade redundante e licenciamento previsível.

VANTAGEM COMPETITIVA SÓ EXISTE QUANDO É ENTREGÁVEL

Ter recurso renovável no país não significa que um terreno específico disponha de potência. O estudo deve chegar ao nível de subestação, fila de conexão, rota de fibra e capacidade de expansão.

A IA cria uma nova geografia de poder computacional

Capacidade de compute vira ativo econômico e estratégico.

Portos, energia, telecomunicações e logística sempre foram infraestruturas decisivas para a competitividade. A IA adiciona uma nova camada: capacidade computacional acelerada. País ou empresa que controla acesso a chips, energia, dados, modelos e conectividade possui maior liberdade para inovar e operar sistemas críticos.

Cinco vetores da nova corrida

- Semicondutores: acesso a GPUs, HBM, packaging avançado e capacidade fabril.
- Energia: disponibilidade de megawatts, preço e prazo de conexão.
- Data centers: densidade, liquid cooling, segurança e operação.
- Modelos e dados: propriedade intelectual, soberania e governança.
- Redes: fibra, DCI, landing stations, IXPs e rotas internacionais diversas.

SOBERANIA DIGITAL NA PRÁTICA

Não significa necessariamente fazer tudo dentro de casa. Significa evitar dependências não gerenciadas, ter alternativas arquiteturais e conhecer onde estão dados, compute, chaves, modelos e rotas.

Como uma empresa deve planejar IA em escala

O roteiro une negócio, dados, segurança e infraestrutura.

Etapa	Pergunta-chave	Entregável
1. Caso de uso	Qual problema e qual KPI?	Business case e critério de sucesso
2. Dados	Onde estão, quem pode acessar?	Mapa de dados e política de uso
3. Modelo	Cloud, API, open model ou privado?	Arquitetura de IA
4. Integração	Quais sistemas e ferramentas?	APIs, RAG, agentes e workflow
5. Risco	O que pode ser automatizado?	Guardrails, aprovação e auditoria
6. Infra	GPU, storage, rede, energia e cooling?	Sizing e desenho técnico
7. Operação	Como observar custo, qualidade e falhas?	SLOs, FinOps e AIOps

DECISÃO BUY VS BUILD

A maior parte das empresas combina modelos via API/cloud com dados e controles próprios. Workloads sensíveis, previsíveis ou intensivos podem justificar infraestrutura privada, colocation ou GPU cloud dedicada. A resposta depende de custo total, latência, soberania e utilização.

A infraestrutura é parte do produto de IA

A próxima fase será definida tanto por software quanto por megawatts, água, fibras e operação.

Em 1956, a inteligência artificial era uma ambição científica. Hoje, ela conversa, cria, analisa e executa. O próximo salto não depende apenas de modelos melhores: depende de capacidade computacional suficiente para atender usuários e empresas com custo, confiabilidade e sustentabilidade.

Por isso, compreender IA exige olhar além do algoritmo. GPUs precisam de memória. GPUs em escala precisam de fabricas de baixa latência. Fabricas precisam de storage. Tudo precisa de energia. Energia vira calor. O calor precisa ser removido. E o resultado precisa atravessar redes nacionais e globais.

A TESE CENTRAL DESTA E-BOOK

A inteligência artificial está se tornando infraestrutura econômica global. Quem planejar compute, energia, cooling e conectividade como um sistema integrado estará melhor preparado para a próxima década.

EnQ DIGITAL

Cloud - IA - Data Center - Conectividade

Infraestrutura Inteligente para Empresas que Não Podem Parar.

Pacote SEO para o site da EnQ Digital

Título SEO sugerido

Inteligência Artificial: história, evolução, GPUs e data centers de IA

Meta description

Entenda a evolução da inteligência artificial, chatbots, imagens e vídeos generativos, impacto econômico, GPUs, liquid cooling, energia e conectividade global de data centers.

Slug sugerido

/inteligencia-artificial-historia-evolucao-data-centers-gpu

Palavras-chave prioritárias

- inteligência artificial
- IA generativa
- infraestrutura de IA
- GPU para inteligência artificial
- data center para IA
- liquid cooling
- Direct Liquid Cooling
- Santos Dumont LNCC
- HPC Brasil
- energia para data center
- data center sustentável
- DCI e DWDM
- cabos submarinos
- AI Factory
- agentes de IA

ESTRATÉGIA DE ARTIGO PILAR

Publique o e-book como download e transforme cada capítulo em um artigo satélite, todos apontando para uma página pilar. Isso cria profundidade semântica em IA + infraestrutura, território diretamente aderente ao posicionamento da EnQ Digital.

Fontes técnicas e institucionais

Dartmouth College - Artificial Intelligence - Our Story. <https://ai.dartmouth.edu/our-story>

Weizenbaum, J. - ELIZA - A Computer Program for the Study of Natural Language Communication. *Communications of the ACM*, 1966

Vaswani et al. - Attention Is All You Need. *NeurIPS*, 2017 - <https://arxiv.org/abs/1706.03762>

UNCTAD - Technology and Innovation Report 2025: Inclusive AI for Development. <https://unctad.org/publication/technology-and-innovation-report-2025>

IEA - Energy and AI (2025). <https://www.iea.org/reports/energy-and-ai>

EPE - Balanço Energético Nacional 2026 - ano-base 2025. <https://www.epe.gov.br/pt/publicacoes-dados-abertos/publicacoes/balanco-energetico-nacional-ben>

LNCC - Configuração do Supercomputador Santos Dumont. <https://sdumont.lncc.br/machine.php>

Eviden/Bull - Santos Dumont, LNCC supercomputer, receives fourfold upgrade. <https://www.bull.com/en/press-releases/santos-dumont-lncc-supercomputer-receives-fourfold-upgrade>

Synergy Research Group - Hyperscale Operators to Account for 67% of all Data Center Capacity by 2031. <https://www.srgresearch.com/articles/hyperscale-operators-to-account-for-67-of-all-data-center-capacity-by-2031>

ITU - Global Connectivity Report 2025. <https://www.itu.int/itu-d/reports/statistics/global-connectivity-report-2025/>

ITU - Submarine Cable Resilience. <https://www.itu.int/digital-resilience/submarine-cables/>

NOTA SOBRE PROJEÇÕES

Números futuros de mercado, consumo de energia e participação de capacidade são cenários ou projeções das fontes indicadas. Eles podem mudar com eficiência, demanda, regulação, cadeia de suprimentos e disponibilidade de energia.

Imagens reais e licenças de uso

Imagem	Crédito	Licença
Computação HPC - NREL	Dennis Schroeder / NREL / U.S. Department of Energy	Domínio público (U.S. Government work)
Conversa ELIZA	Autor desconhecido / Wikimedia Commons	Texto representado em domínio público conforme página do arquivo
Santos Dumont	Neila Rocha/MCTIC / Ministério da Ciência, Tecnologia e Inovações	CC BY 2.0
Cooling de data center	Rsparks3	CC0 1.0
Energia solar	Dennis Schroeder / U.S. Department of Energy	Domínio público (U.S. Government work)
Energia eólica	U.S. Department of Energy	Domínio público (U.S. Government work)
Fibra óptica	MTA Capital Construction Mega Projects	CC BY 2.0
Cabo submarino	~~ ItsJegh ~~	CC BY 2.0

Observação editorial

As fotografias foram selecionadas por disponibilidade pública e licenças que permitem reutilização, respeitando os créditos indicados. Imagens CC BY 2.0 exigem atribuição. Antes da publicação final em outros formatos, preserve esta página de créditos ou os créditos junto das imagens.

Wikimedia Commons: High performance computing (HPC) data center (9091835239).jpg

Wikimedia Commons: ELIZA conversation.png

Wikimedia Commons: Supercomputador Santos Dumont (49128915827).jpg

Wikimedia Commons: Data center cooling tower.jpg

Wikimedia Commons: Installing Solar Panels (7336033672).jpg

Wikimedia Commons: Power County Wind Farm 003.jpg

Wikimedia Commons: Layout and installation of fiber optic cables... (49248066228).jpg

Wikimedia Commons: Submarine Cable (6177266137).jpg