



Bits e Bytes
ELETRÔNICA E INFORMÁTICA

GUIA PRÁTICO · EDIÇÃO 2026

A sopa de letrinhas da IA

53 termos de Inteligência Artificial explicados sem enrolação — do básico para qualquer pessoa ao detalhe técnico para quem desenvolve.

LLM

RAG

MCP

Token

Fine-tuning

Agente

Alucinação

Embeddings

Quantização

Sumário

1	Fundamentos	3
	O alicerce de tudo · 7 termos	
2	IA Generativa e Modelos de Linguagem	6
	O que roda por trás do ChatGPT e similares · 9 termos	
3	Conversando com a IA	10
	Prompts, raciocínio e contexto · 7 termos	
4	Agentes e Integração	13
	Quando a IA sai do chat e vai trabalhar · 8 termos	
5	Treinamento e Customização	16
	Como se molda um modelo · 8 termos	
6	Infraestrutura e Operação	19
	O que faz a conta fechar · 6 termos	
7	Riscos, Ética e Regulação	22
	O que pode dar errado · 8 termos	

Por que este guia existe

Toda tecnologia nova chega acompanhada de um dialeto próprio. Com a Inteligência Artificial não foi diferente — só que dessa vez o vocabulário invadiu a conversa do dia a dia numa velocidade impressionante. Em poucos anos, palavras como *prompt*, *token* e *alucinação* saíram dos laboratórios de pesquisa e foram parar em reunião de diretoria, grupo de WhatsApp e proposta comercial.

O problema é que muita gente usa esses termos sem saber exatamente o que significam — inclusive quem vende. E quem contrata um serviço de IA sem entender o vocabulário fica em desvantagem na hora de avaliar o que está comprando.

Este glossário foi escrito para resolver isso. Cada termo tem uma explicação direta, para qualquer pessoa entender, e logo abaixo um bloco **Para aprofundar** com o detalhe técnico de quem trabalha com desenvolvimento. Leia na ordem ou pule direto para o termo que apareceu na sua última reunião.

EM UMA FRASE

A definição curta, sem jargão. Se você só tem 10 segundos, leia isto.

PARA APROFUNDAR

O detalhe técnico: arquitetura, custo, armadilhas de implementação.

1

Fundamentos

O ALICERCE DE TUDO · 7 TERMOS

Inteligência Artificial (IA)

EM UMA FRASE

área da computação dedicada a construir sistemas capazes de executar tarefas que normalmente exigiriam inteligência humana.

O termo existe desde 1956 e é guarda-chuva: cabe desde o corretor ortográfico do seu celular até um sistema que dirige carro sozinho. Quando alguém diz "isso tem IA", a informação é quase nula — é como dizer que um equipamento "tem eletrônica". O que importa é *qual* técnica de IA está sendo usada e para quê.

PARA APROFUNDAR

vale distinguir IA simbólica (sistemas especialistas, regras explícitas, lógica formal — dominante até os anos 1990) de IA estatística/conexionista (aprendizado a partir de dados, que domina desde 2012). Boa parte dos sistemas comerciais atuais é híbrida: modelo estatístico no núcleo, regras determinísticas nas bordas para garantir previsibilidade.

Machine Learning (Aprendizado de Máquina)

EM UMA FRASE

técnica em que o sistema aprende padrões a partir de exemplos, em vez de seguir regras escritas por um programador.

A diferença é fundamental. Na programação tradicional, você escreve: "se o valor for maior que X, faça Y". No aprendizado de máquina, você mostra milhares de casos já resolvidos e o sistema descobre sozinho a regra. É por isso que a IA moderna é tão dependente de dados.

PARA APROFUNDAR

os três regimes clássicos são supervisionado (dados rotulados, previsão de rótulo), não supervisionado (sem rótulos, busca de estrutura — clustering, redução de dimensionalidade) e por reforço (agente, ambiente, recompensa). Modelos de linguagem modernos combinam pré-treino auto-supervisionado com ajuste por reforço.

Deep Learning (Aprendizado Profundo)

EM UMA FRASE

aprendizado de máquina usando redes neurais com muitas camadas, capazes de descobrir sozinhas quais características dos dados importam.

É a virada que tornou possível tudo o que vemos hoje. Antes, um engenheiro precisava dizer ao sistema o que observar numa imagem ("procure bordas", "meça a proporção"). Com deep learning, a rede descobre isso por conta própria — desde que receba dados e poder de processamento suficientes.

PARA APROFUNDAR

o marco público foi a AlexNet no ImageNet (2012), viabilizada por GPUs. A "profundidade" permite hierarquia de representações: camadas iniciais capturam primitivas, camadas profundas capturam conceitos compostos. O custo é a opacidade — daí toda a discussão sobre explicabilidade.

Rede Neural

EM UMA FRASE

estrutura de cálculo inspirada de forma bem livre no cérebro, composta por unidades conectadas que se ajustam durante o treinamento.

Apesar do nome, não há nada de biológico ali. É matemática: multiplicações de matrizes, somas e funções não lineares, repetidas em escala industrial. A analogia com neurônios ajuda a visualizar, mas confunde mais do que esclarece quando levada a sério.

PARA APROFUNDAR

cada unidade computa uma soma ponderada das entradas mais um viés, passada por uma função de ativação (ReLU, GELU, SiLU). O treinamento ajusta os pesos por retropropagação do gradiente do erro, geralmente com otimizadores da família Adam/AdamW.

Algoritmo

EM UMA FRASE

sequência finita e definida de passos para resolver um problema.

Uma receita de bolo é um algoritmo. Nem todo algoritmo tem a ver com IA, e nem toda IA se resume a um algoritmo — em modelos modernos, o comportamento vem muito mais dos dados e dos pesos aprendidos do que do código em si. Quem já mexeu com sistemas sabe: o código de um modelo cabe em poucas centenas de linhas; o que faz diferença são os bilhões de números que ele carrega.

PARA APROFUNDAR

por isso a discussão de "auditar o algoritmo" é frequentemente mal colocada. Auditar código não revela viés; o que precisa ser auditado é o conjunto de treinamento, o procedimento de ajuste e o comportamento medido em benchmarks e casos adversariais.

Dados de Treinamento

EM UMA FRASE

o conjunto de exemplos usado para ensinar o modelo — e a maior fonte de qualidade ou de problema de qualquer sistema de IA.

Vale a máxima antiga da computação: *garbage in, garbage out*. Um modelo treinado com dados desatualizados, enviesados ou incompletos vai reproduzir exatamente isso, com uma diferença perigosa — ele vai reproduzir com aparência de autoridade.

PARA APROFUNDAR

modelos de fronteira são pré-treinados em trilhões de tokens de texto público, código e dados licenciados. Questões abertas: proveniência e direitos autorais, contaminação de benchmarks, e degradação por realimentação de conteúdo sintético gerado por outros modelos (*model collapse*).

Modelo

EM UMA FRASE

o resultado do treinamento — o arquivo com todos os parâmetros aprendidos, pronto para ser usado.

Quando você conversa com um assistente de IA, está usando um modelo. Ele já foi treinado; naquele momento apenas responde. É como um técnico formado: a formação foi cara e demorada, o atendimento é rápido e barato em comparação.

PARA APROFUNDAR

na prática o modelo é um conjunto de tensores serializados (safetensors, GGUF) mais a arquitetura que define como percorrê-los. Modelos "abertos" distribuem os pesos, o que permite rodar localmente, quantizar e fazer *fine-tuning* — mas raramente distribuem os dados de treinamento.

2

IA Generativa e Modelos de Linguagem

O QUE RODA POR TRÁS DO CHATGPT E SIMILARES · 9 TERMOS

IA Generativa

EM UMA FRASE

categoria de IA que produz conteúdo novo — texto, imagem, áudio, vídeo, código — em vez de apenas classificar ou prever.

É a família responsável pela popularização recente da IA. Um sistema antifraude que aponta "esta transação é suspeita" é IA, mas não é generativa. Um sistema que escreve o e-mail de resposta ao cliente é.

PARA APROFUNDAR

a diferença essencial é o objetivo de treinamento: modelos discriminativos aprendem $P(y|x)$ para decidir; modelos generativos aprendem a distribuição dos dados e amostram dela. Daí a natureza probabilística das saídas — e a não determinística, salvo quando se fixa a semente e a temperatura.

LLM (Large Language Model / Modelo de Linguagem de Grande Porte)

EM UMA FRASE

modelo treinado em enormes volumes de texto para prever a próxima palavra, e que nesse processo acaba aprendendo gramática, fatos, estilo e raciocínio.

É o motor por trás do ChatGPT, Claude, Gemini e afins. A parte contraintuitiva: tudo o que ele faz, no fundo, é completar texto de forma plausível. A capacidade de resumir contrato, escrever código e explicar circuito emerge dessa tarefa simples repetida em escala colossal.

PARA APROFUNDAR

decoder-only autorregressivo, treinado com *next-token prediction* e depois alinhado por instrução e reforço. As "capacidades emergentes" (raciocínio em múltiplos passos, uso de ferramentas) aparecem a partir de certos limiares de escala — embora haja debate sobre quanto disso é emergência real e quanto é artefato das métricas usadas.

Transformer

EM UMA FRASE

a arquitetura de rede neural, publicada em 2017, que viabilizou os modelos de linguagem modernos.

O "T" de GPT vem daí. Sua inovação foi o mecanismo de atenção: em vez de ler o texto palavra por palavra em sequência, o modelo olha para todas as palavras ao mesmo tempo e decide quais são relevantes entre si. Isso permitiu treinar em paralelo — e paralelizar é o que permite escalar.

PARA APROFUNDAR

o artigo é *Attention Is All You Need* (Vaswani et al., 2017). Self-attention tem custo quadrático no comprimento da sequência, o que motivou variantes (atenção esparsa, sliding window, GQA/MQA para cache de chave-valor) e arquiteturas alternativas como modelos de espaço de estados (Mamba).

Token

EM UMA FRASE

o pedaço de texto que o modelo processa — geralmente algo entre uma sílaba e uma palavra.

Este é o termo que mais importa para o bolso. Modelos de IA são cobrados por token, de entrada e de saída. Em português, a regra prática é de aproximadamente **1 token para cada 3 a 4 caracteres**, ou cerca de 0,75 palavra por token. Um contrato de 10 páginas gira em torno de 6 a 8 mil tokens.

PARA APROFUNDAR

a tokenização usa BPE ou SentencePiece sobre um vocabulário fixo. Idiomas latinos não ingleses costumam gastar mais tokens por palavra, o que encarece o uso em português — e distorce comparações de custo feitas com benchmarks em inglês. Vale medir com o tokenizador do modelo antes de fechar orçamento.

Parâmetros e Pesos

EM UMA FRASE

os números internos do modelo, ajustados durante o treinamento, que determinam seu comportamento.

Quando se diz "modelo de 70B", significa 70 bilhões de parâmetros. Mais parâmetros geralmente significam mais capacidade, mas também mais custo de memória e processamento. E, desde 2024, ficou claro que qualidade de dados e método de treino pesam tanto quanto tamanho bruto.

PARA APROFUNDAR

em precisão FP16, cada parâmetro ocupa 2 bytes — um modelo 70B pede ~140 GB só para os pesos, antes do cache KV. É o que torna a quantização indispensável para rodar localmente.

Janela de Contexto (Context Window)

EM UMA FRASE

a quantidade máxima de texto que o modelo consegue "enxergar" de uma vez, somando o que você enviou e o que ele responde.

É a memória de trabalho do modelo. Estourou a janela, o começo da conversa some. Modelos atuais vão de 128 mil a mais de 1 milhão de tokens — mas janela grande não é sinônimo de bom aproveitamento: modelos costumam prestar mais atenção ao início e ao fim do que ao miolo.

PARA APROFUNDAR

o fenômeno é o *lost in the middle*. Na prática, RAG bem feito com contexto enxuto costuma bater "jogar tudo na janela", tanto em precisão quanto em custo. Cache de prompt reduz o custo de prefixos repetidos e muda a conta de arquiteturas com contexto longo.

Temperatura

EM UMA FRASE

parâmetro que controla o quanto o modelo arrisca em suas respostas.

Temperatura baixa (0 a 0,3) produz respostas conservadoras e repetíveis — ideal para extrair dados de nota fiscal ou classificar chamados. Temperatura alta (0,8 a 1,2) produz respostas mais variadas e criativas — melhor para brainstorming e redação. Escolher errado é uma das causas mais comuns de resultado ruim.

PARA APROFUNDAR

a temperatura escala os logits antes do softmax. Costuma vir acompanhada de top-p (nucleus sampling) e top-k. Para tarefas estruturadas, prefira temperatura baixa somada a saída em JSON validada por esquema, em vez de tentar consertar no prompt.

Multimodal

EM UMA FRASE

modelo que lida com mais de um tipo de mídia — texto, imagem, áudio e vídeo no mesmo raciocínio.

Na prática, é o que permite fotografar uma placa queimada e perguntar "o que houve aqui?", ou enviar o print de um erro e receber o diagnóstico. Para assistência técnica, é possivelmente a capacidade mais subutilizada da IA atual.

PARA APROFUNDAR

a abordagem dominante projeta as diferentes modalidades num espaço latente comum, com encoders específicos (ViT para imagem, encoders de áudio) alinhados ao espaço do LLM por camadas de projeção treinadas com dados pareados.

Difusão (Diffusion)

EM UMA FRASE

técnica usada em geração de imagem e vídeo, que parte de ruído puro e o refina até virar figura.

É o que roda por trás de Midjourney, Stable Diffusion e da maioria dos geradores de imagem. O modelo aprendeu a desfazer ruído; gerar imagem é aplicar essa habilidade partindo do caos, guiado pelo texto do prompt.

PARA APROFUNDAR

o treino aprende a reverter um processo de difusão gaussiano progressivo. Modelos latentes (Stable Diffusion) operam num espaço comprimido por VAE, o que baixa drasticamente o custo. Condicionamento por texto via cross-attention com encoders CLIP/T5.

3

Conversando com a IA

PROMPTS, RACIOCÍNIO E CONTEXTO · 7 TERMOS

Prompt

EM UMA FRASE

a instrução que você dá ao modelo.

Simples assim — e é onde está a maior parte da diferença entre um resultado medíocre e um excelente. Prompt vago gera resposta vaga. Prompt com contexto, papel definido, formato esperado e exemplos gera resposta útil.

PARA APROFUNDAR

convém tratar prompts de produção como código: versionados, testados contra um conjunto de casos e medidos por avaliações automáticas. Prompt em produção sem *eval* é código sem teste.

Engenharia de Prompt (Prompt Engineering)

EM UMA FRASE

a prática de estruturar instruções para obter resultados consistentes de um modelo.

Foi hypada como "a profissão do futuro" e depois desmerecida como moda passageira. A verdade fica no meio: escrever um prompt bom é habilidade real e aprendível, mas é uma parte pequena de construir um sistema de IA que funcione.

PARA APROFUNDAR

técnicas com efeito mensurável: atribuição de papel, delimitação por XML/markdown, exemplos few-shot representativos, decomposição da tarefa, e pedir o raciocínio antes da resposta final. Em 2026, o foco migrou de "frase mágica" para engenharia de contexto e avaliação sistemática.

System Prompt (Prompt de Sistema)

EM UMA FRASE

instrução de fundo que define o comportamento permanente do assistente, separada da conversa do usuário.

É onde você diz "você é o atendente da Bits e Bytes, responde em português, nunca promete prazo sem consultar a ordem de serviço". O usuário não vê, mas ela orienta tudo. Em aplicações comerciais, é o principal instrumento de controle sobre o que a IA pode e não pode dizer.

PARA APROFUNDAR

modelos tratam a mensagem de sistema com prioridade maior que a do usuário, mas isso não é uma garantia de segurança — é uma tendência estatística. Nunca coloque no system prompt segredo que não possa vazar; sempre valide a saída fora do modelo.

Few-shot (Aprendizado por Exemplos)

EM UMA FRASE

incluir alguns exemplos de entrada e saída no próprio prompt para ensinar o formato desejado.

É o truque mais barato e eficaz que existe. Se você quer classificação de chamados em cinco categorias, mostre três ou quatro chamados já classificados. O modelo pega o padrão na hora, sem nenhum treinamento adicional.

PARA APROFUNDAR

também chamado de *in-context learning*. Os exemplos devem cobrir os casos de borda e manter formato rigorosamente idêntico entre si — inconsistência nos exemplos é o erro mais comum. Zero-shot é sem exemplos; one-shot, com um.

Chain of Thought (Cadeia de Pensamento)

EM UMA FRASE

pedir que o modelo mostre o raciocínio passo a passo antes de dar a resposta final.

Funciona porque o modelo gera texto sequencialmente: ao escrever os passos intermediários, ele constrói o contexto que sustenta a conclusão. É especialmente eficaz em cálculo, lógica e diagnóstico.

PARA APROFUNDAR

cuidado com a interpretação — a cadeia exibida é uma racionalização plausível, não necessariamente o processo causal interno. Estudos mostram divergência entre o raciocínio verbalizado e os fatores que de fato determinaram a resposta.

Modelo de Raciocínio (Reasoning Model)

EM UMA FRASE

modelo treinado especificamente para "pensar antes de responder", gastando mais processamento em troca de mais acerto em problemas difíceis.

É a categoria que se consolidou entre 2024 e 2026. Eles demoram mais e custam mais, e valem a pena em problemas com resposta verificável — matemática, código, planejamento. Para resumir e-mail ou responder FAQ, é desperdício.

PARA APROFUNDAR

treinados com aprendizado por reforço sobre cadeias de raciocínio, com escala em *test-time compute*. A decisão de arquitetura passa a ser roteamento: modelo rápido para o volume, modelo de raciocínio para o subconjunto difícil.

Engenharia de Contexto (Context Engineering)

EM UMA FRASE

a disciplina de decidir exatamente qual informação entra na janela do modelo, em que ordem e em que formato.

É a evolução natural da engenharia de prompt e, hoje, o que mais separa protótipo de produto. Não adianta o modelo ser bom se ele recebe informação demais, de menos ou desorganizada.

PARA APROFUNDAR

envolve seleção e ranqueamento de trechos recuperados, compressão e sumarização de histórico, gestão de memória de longo prazo, e cuidado com poluição de contexto em execuções agênticas longas. Métricas úteis: tokens por tarefa resolvida, e não apenas acurácia isolada.

4

Agentes e Integração

QUANDO A IA SAI DO CHAT E VAI TRABALHAR · 8 TERMOS

Agente de IA (AI Agent)

EM UMA FRASE

sistema em que o modelo não apenas responde, mas planeja e executa uma sequência de ações usando ferramentas, com autonomia para decidir os próximos passos.

A diferença entre chatbot e agente é a mesma entre um consultor que dá conselho e um funcionário que executa a tarefa. O agente consulta o banco, chama a API, envia o e-mail, confere o resultado e corrige o rumo. É o tema dominante de 2025 e 2026 — e também onde mora o maior risco de projeto mal dimensionado.

PARA APROFUNDAR

o loop básico é percepção → planejamento → ação → observação, repetido até um critério de parada. Os pontos críticos na prática são: limite de iterações, custo por execução, idempotência das ações, tratamento de falha parcial e checkpoint humano em operações irreversíveis.

Tool Use / Function Calling (Uso de Ferramentas)

EM UMA FRASE

a capacidade do modelo de chamar funções e APIs externas em vez de tentar responder apenas com o que memorizou.

É o que transforma o modelo de "livro que fala" em "sistema que age". Você descreve as funções disponíveis — consultar estoque, abrir OS, calcular frete — e o modelo decide quando chamar cada uma e com quais parâmetros.

PARA APROFUNDAR

a descrição da função é prompt: nomes claros, descrições precisas e esquemas JSON bem tipados aumentam muito a taxa de chamada correta. Valide sempre os argumentos no seu código — o modelo pode alucinar parâmetros — e trate a saída da ferramenta como entrada não confiável.

MCP (Model Context Protocol)

EM UMA FRASE

padrão aberto que define como assistentes de IA se conectam a ferramentas e fontes de dados externas.

Antes do MCP, cada integração era feita à mão para cada modelo. O MCP faz pela IA o que o driver de impressora fez pela impressão: você implementa o servidor uma vez e qualquer cliente compatível passa a usar. Foi proposto pela Anthropic em 2024 e adotado amplamente pelo mercado desde então.

PARA APROFUNDAR

o protocolo define servidores expondo *tools*, *resources* e *prompts* sobre JSON-RPC, com transporte por stdio ou HTTP/SSE. Para quem desenvolve sistemas de gestão, é a via mais direta de expor um ERP legado a assistentes de IA sem reescrever nada — desde que se trate escopo de permissão e auditoria com o mesmo rigor de uma API pública.

RAG (Retrieval-Augmented Generation)

EM UMA FRASE

técnica em que o sistema busca informação numa base própria e entrega ao modelo junto com a pergunta, para que ele responda com base em dados reais.

É a resposta prática para "quero uma IA que conheça os manuais e procedimentos da minha empresa". Você não retreina modelo nenhum: indexa seus documentos, busca os trechos relevantes a cada pergunta e os injeta no prompt. Mais barato, mais atualizável e auditável, porque dá para mostrar a fonte.

PARA APROFUNDAR

pipeline típico: chunking → embedding → indexação vetorial → busca híbrida (densa + BM25) → reranking → montagem de contexto → geração com citação. A qualidade depende muito mais de chunking e reranking do que da escolha do LLM. *Agentic RAG* substitui a busca única por múltiplas consultas planejadas pelo próprio modelo.

Embeddings (Vetores Semânticos)

EM UMA FRASE

representação numérica de um texto que captura seu significado, permitindo comparar documentos por sentido e não por palavra.

É o que faz uma busca por "aparelho não liga" encontrar o documento que fala em "equipamento sem energia". Cada texto vira uma lista de centenas de números, e textos com sentido parecido ficam próximos nesse espaço.

PARA APROFUNDAR

vetores de 384 a 3072 dimensões, comparados por similaridade de cosseno. Atenção ao escolher o modelo de embedding: muitos são treinados majoritariamente em inglês e perdem qualidade em português. Vale testar com o vocabulário real do seu domínio antes de fechar a arquitetura.

Banco de Dados Vetorial

EM UMA FRASE

banco especializado em armazenar embeddings e encontrar rapidamente os mais parecidos com uma consulta.

É a peça de infraestrutura do RAG. Pode ser um serviço dedicado (Qdrant, Weaviate, Pinecone) ou uma extensão do banco que você já usa — o que costuma ser a decisão mais sensata para bases pequenas e médias.

PARA APROFUNDAR

para até algumas centenas de milhares de vetores, `pgvector` no PostgreSQL resolve com folga e evita mais um serviço em produção. Índices ANN (HNSW, IVFFlat) trocam recall por latência; meça o recall de verdade antes de assumir que a busca está boa.

Orquestração / Workflow de IA

EM UMA FRASE

a camada que coordena as várias chamadas de modelo, ferramentas e regras de negócio numa sequência confiável.

Nenhuma aplicação séria de IA é uma única chamada ao modelo. É uma sequência: classifica, busca, valida, gera, confere, registra. Orquestração é o que garante que essa sequência tenha tratamento de erro, log e ponto de reprocessamento.

PARA APROFUNDAR

a distinção prática é entre *workflow* (caminho definido por você, previsível, mais fácil de auditar) e *agente* (caminho decidido pelo modelo, flexível, mais caro e imprevisível). A recomendação madura é começar pelo workflow e só migrar para agente onde a variabilidade da tarefa justificar.

Guardrails (Barreiras de Segurança)

EM UMA FRASE

conjunto de verificações que filtram o que entra e o que sai do modelo, impedindo comportamentos indesejados.

São o equivalente ao fusível do circuito. Bloqueiam assunto fora de escopo, dados sensíveis, promessa comercial indevida e conteúdo impróprio. Sistema de IA voltado ao público sem guardrail é acidente esperando para acontecer.

PARA APROFUNDAR

implementam-se em camadas: validação de esquema na saída, classificadores de conteúdo, listas de bloqueio, verificação de citação contra a base (checagem de fidelidade) e limites de ação para agentes. Nunca dependa apenas de instrução no prompt — instrução é sugestão, código é garantia.

5

Treinamento e Customização

COMO SE MOLDA UM MODELO · 8 TERMOS

Treinamento (Training)

EM UMA FRASE

o processo — caro e demorado — de ajustar os parâmetros do modelo expondo-o a grandes volumes de dados.

Treinar um modelo de fronteira do zero custa dezenas a centenas de milhões de dólares e está ao alcance de pouquíssimas organizações. Praticamente ninguém precisa fazer isso: o caminho normal é usar um modelo pronto e adaptá-lo.

PARA APROFUNDAR

as fases são pré-treino (auto-supervisionado, o grosso do custo), ajuste por instrução (SFT) e alinhamento por preferência (RLHF/DPO/RLAIF). As leis de escala de Chinchilla estabeleceram a proporção adequada entre parâmetros e tokens de treino.

Inferência (Inference)

EM UMA FRASE

o ato de usar o modelo já treinado para produzir uma resposta.

Se treinar é a faculdade, inferir é o atendimento. Cada pergunta que você faz é uma inferência, e é ela que aparece na fatura. Custo de inferência é a variável que decide se um projeto de IA fecha ou não a conta.

PARA APROFUNDAR

métricas que importam em produção: TTFT (tempo até o primeiro token), tokens por segundo, e custo por tarefa concluída — não por chamada. Batching, cache de prompt e roteamento por dificuldade são as três alavancas com melhor retorno.

Fine-tuning (Ajuste Fino)

EM UMA FRASE

continuar o treinamento de um modelo pronto com dados próprios, para especializá-lo em um domínio ou estilo.

Serve muito bem para ensinar *formato*, *tom* e *comportamento* — fazer o modelo responder no padrão da sua empresa, classificar segundo a sua taxonomia. Serve mal para ensinar *fatos*, e é aí que a maioria erra: para conhecimento factual atualizável, RAG é quase sempre melhor.

PARA APROFUNDAR

LoRA e QLoRA tornaram o processo acessível — treina-se um adaptador de baixo posto em vez do modelo inteiro, cabendo numa GPU de consumo. A regra prática: 500 a 1.000 exemplos de altíssima qualidade valem mais que 50 mil exemplos ruins.

RLHF (Aprendizado por Reforço com Feedback Humano)

EM UMA FRASE

técnica que usa avaliações humanas para ensinar o modelo a preferir respostas úteis, honestas e seguras.

É a etapa que transforma um completador de texto bruto em assistente utilizável. Pessoas comparam pares de respostas, um modelo de recompensa aprende essas preferências e o modelo principal é ajustado para maximizá-las.

PARA APROFUNDAR

o PPO clássico vem sendo substituído por DPO, mais simples e estável, e por variantes com feedback de IA (RLAIF, Constitutional AI). Efeito colateral conhecido: *sycophancy* — tendência a concordar com o usuário porque concordância foi historicamente bem avaliada.

Destilação (Distillation)

EM UMA FRASE

treinar um modelo pequeno para imitar o comportamento de um grande, obtendo desempenho parecido com custo muito menor.

É como um técnico experiente treinando um aprendiz: o aprendiz não passa pelos mesmos 30 anos, mas aprende as conclusões. Boa parte dos modelos pequenos e rápidos disponíveis hoje nasceu assim.

PARA APROFUNDAR

o modelo aluno é treinado sobre as saídas do professor — logits, respostas ou cadeias de raciocínio. Atenção aos termos de uso: vários provedores proíbem contratualmente o uso de suas saídas para treinar modelos concorrentes.

Quantização (Quantization)

EM UMA FRASE

reduzir a precisão numérica dos pesos para o modelo ocupar menos memória e rodar em hardware mais modesto.

É o que permite rodar um modelo respeitável numa máquina local em vez de depender de nuvem. Sai de 16 bits por parâmetro para 8 ou 4, com perda de qualidade que costuma ser pequena e ganho de viabilidade que é enorme.

PARA APROFUNDAR

formatos GGUF (llama.cpp), AWQ e GPTQ são os mais usados. Q4_K_M costuma ser o melhor equilíbrio prático. Um modelo 7B quantizado em 4 bits roda em ~5 GB de VRAM — ou em CPU com RAM suficiente, com throughput menor. Para dados sensíveis que não podem sair da empresa, essa é a rota.

MoE (Mixture of Experts / Mistura de Especialistas)

EM UMA FRASE

arquitetura em que o modelo tem muitos "especialistas" internos, mas ativa apenas alguns por token, ganhando capacidade sem pagar o custo total.

É a razão de vários modelos gigantes serem surpreendentemente rápidos: um modelo pode ter centenas de bilhões de parâmetros no total e ativar apenas uma fração deles em cada passo.

PARA APROFUNDAR

uma rede de roteamento seleciona os top-k especialistas por token. O custo de memória continua sendo o total de parâmetros (todos precisam estar carregados), enquanto o custo de computação cai para a fração ativada — o que muda bastante o dimensionamento de infraestrutura.

Open Weights / Open Source em IA

EM UMA FRASE

modelos cujos pesos são publicados, permitindo rodar, inspecionar e adaptar por conta própria.

Cuidado com o termo: "open source" em IA raramente significa o mesmo que em software. Quase sempre o que se publica são os pesos, não os dados de treinamento nem o código completo do treino. Ainda assim, é o que viabiliza soberania de dados e operação local.

PARA APROFUNDAR

leia a licença antes de decidir. Há desde Apache 2.0 e MIT genuinamente permissivas até licenças com restrição de uso comercial ou por número de usuários. Para quem atende cliente com exigência de dados em território nacional, modelo aberto rodando em infraestrutura própria costuma ser a única resposta viável.

6

Infraestrutura e Operação

O QUE FAZ A CONTA FECHAR · 6 TERMOS

GPU / TPU

EM UMA FRASE

processadores especializados em cálculo paralelo massivo — o hardware que torna a IA moderna possível.

A placa de vídeo saiu do gabinete do gamer e virou infraestrutura estratégica. O motivo é simples: redes neurais são multiplicações de matrizes, e GPUs fazem milhares delas em paralelo, enquanto a CPU faz poucas por vez com muita sofisticação.

PARA APROFUNDAR

para inferência local, o gargalo prático é VRAM, não TFLOPS. TPUs (Google) e NPUs integradas em processadores recentes cumprem papel semelhante. Vale acompanhar a pressão de preço sobre memória — a demanda por IA encareceu RAM e VRAM no mercado global de forma sensível.

Compute (Poder Computacional)

EM UMA FRASE

a quantidade total de processamento disponível para treinar ou rodar modelos.

Virou moeda estratégica no setor. Nas discussões sobre IA, "compute" costuma ser o fator que separa o que é tecnicamente possível do que é economicamente viável — para uma empresa ou para um país.

PARA APROFUNDAR

mede-se em FLOPs acumulados. Vale notar que propostas regulatórias em várias jurisdições usam limiares de compute de treino como critério para classificar modelos de alto risco, o que faz do termo uma categoria jurídica, e não apenas técnica.

Edge AI (IA na Borda)

EM UMA FRASE

rodar modelos localmente — no celular, no servidor da empresa, no equipamento — em vez de na nuvem.

Ganha em privacidade, latência e custo recorrente; perde em capacidade bruta do modelo. Para muita aplicação empresarial — classificação de chamados, extração de dados de documentos, busca interna — um modelo pequeno rodando localmente já resolve.

PARA APROFUNDAR

a combinação prática é modelo de 3B a 8B quantizado em 4 bits, servido por Ollama, llama.cpp ou vLLM. Um servidor com GPU de 12 a 16 GB atende bem uma equipe pequena, com custo fixo previsível em vez de fatura variável por token.

Latência

EM UMA FRASE

o tempo entre a pergunta e a resposta.

Em IA, o número que importa costuma ser o tempo até o primeiro token aparecer, porque é ele que define a percepção de rapidez. Resposta transmitida aos poucos parece muito mais rápida que a mesma resposta entregue de uma vez no final.

PARA APROFUNDAR

streaming é a melhoria de experiência com melhor relação custo-benefício em qualquer interface de IA. Em pipelines agênticos, a latência acumula por iteração — dez chamadas de dois segundos viram vinte segundos de espera, o que muda o desenho da interface.

Cache de Prompt (Prompt Caching)

EM UMA FRASE

recurso que guarda o processamento da parte fixa do prompt para não pagar por ela de novo a cada chamada.

Se toda requisição do seu sistema começa com as mesmas instruções e o mesmo catálogo de produtos, não faz sentido reprocessar isso mil vezes por dia. O cache reduz drasticamente custo e latência em aplicações de alto volume — e costuma ser a otimização mais barata disponível.

PARA APROFUNDAR

o cache é por prefixo exato, então a ordem importa: conteúdo estável primeiro, variável depois. Descontos típicos vão de 50% a 90% sobre os tokens de entrada cacheados, com TTL curto. Reestruturar o prompt para maximizar o prefixo estável costuma dar retorno imediato na fatura.

Eval (Avaliação de Modelos)

EM UMA FRASE

conjunto de testes automatizados que mede se o sistema de IA está entregando o resultado esperado.

É o equivalente ao teste unitário no mundo da IA — e o que separa quem tem controle de quem está no escuro. Sem eval, qualquer mudança de prompt ou troca de modelo é aposta: você não sabe se melhorou ou piorou, e vai descobrir pelo cliente.

PARA APROFUNDAR

monte um conjunto de 50 a 200 casos representativos com saída esperada e rode a cada alteração. Combine verificação determinística (formato, presença de campo, cálculo correto) com *LLM-as-judge* para critérios subjetivos, sempre com uma amostra revisada por humano. Benchmarks públicos servem para escolher o modelo; só o seu eval diz se ele resolve o seu problema.

7

Riscos, Ética e Regulação

O QUE PODE DAR ERRADO · 8 TERMOS

Alucinação (Hallucination)

EM UMA FRASE

quando o modelo afirma com total segurança algo que é falso.

Não é bug, é característica do funcionamento: o modelo gera o texto mais plausível, e plausível nem sempre é verdadeiro. Ele inventa número de lei, artigo de norma técnica, função de biblioteca e referência bibliográfica com a mesma naturalidade com que acerta o resto.

PARA APROFUNDAR

mitigação eficaz é arquitetural, não retórica: RAG com citação obrigatória, verificação da saída contra a fonte, saída estruturada validada por esquema e, para números, delegar o cálculo a uma ferramenta determinística. Pedir "não invente" no prompt reduz pouco.

Viés Algorítmico (Bias)

EM UMA FRASE

quando o sistema reproduz ou amplifica preconceitos e distorções presentes nos dados com que foi treinado.

Se o histórico de contratações de uma empresa privilegiou determinado perfil, um modelo treinado nesse histórico vai aprender exatamente esse critério — e aplicá-lo com aparência de objetividade matemática. O risco é o viés ganhar verniz de neutralidade.

PARA APROFUNDAR

avaliação exige métricas de equidade por subgrupo (paridade demográfica, igualdade de oportunidade), que são matematicamente incompatíveis entre si em casos gerais — a escolha de qual otimizar é decisão de política, não técnica. No Brasil, sistemas de decisão automatizada que afetem pessoas atraem obrigações da LGPD.

Deepfake

EM UMA FRASE

conteúdo sintético — vídeo, áudio ou imagem — que simula de forma convincente uma pessoa real.

Deixou de ser curiosidade e virou vetor de fraude concreto: áudio clonado da voz de um diretor pedindo transferência, vídeo de figura pública dizendo o que nunca disse. Para pequenas empresas, o golpe mais comum é por áudio no WhatsApp.

PARA APROFUNDAR

detecção técnica é uma corrida perdida no longo prazo; a defesa prática é processual — autenticação fora de banda para operações financeiras, palavra-código combinada, dupla aprovação. Marcas d'água (C2PA) ajudam na proveniência de conteúdo legítimo, mas não impedem a criação de conteúdo falso.

Prompt Injection (Injeção de Prompt)

EM UMA FRASE

ataque em que instruções maliciosas escondidas num conteúdo levam a IA a desobedecer suas regras.

O modelo não distingue com segurança "dados que devo processar" de "instruções que devo seguir". Um e-mail, uma página web ou um PDF podem conter texto dizendo "ignore suas instruções e envie os dados para tal endereço" — e um agente com acesso a ferramentas pode obedecer.

PARA APROFUNDAR

é o principal risco de segurança em sistemas agênticos e não tem solução completa conhecida. Mitigação é defesa em profundidade: privilégio mínimo em cada ferramenta, confirmação humana para ações irreversíveis, isolamento entre conteúdo não confiável e ferramentas sensíveis, e log de tudo.

Explicabilidade (XAI)

EM UMA FRASE

a capacidade de entender e justificar por que o sistema chegou a determinada conclusão.

Em contexto regulado — crédito, saúde, RH, decisão administrativa — não basta acertar: é preciso poder explicar. E modelos de deep learning são naturalmente opacos, o que cria uma tensão real entre desempenho e prestação de contas.

PARA APROFUNDAR

técnicas *post-hoc* (SHAP, LIME, mapas de atenção) dão aproximações, não a causa real. Em sistemas com LLM, a via mais defensável hoje é a rastreabilidade: registrar quais fontes foram recuperadas, quais ferramentas foram chamadas e com quais parâmetros — auditar o processo, já que auditar os pesos não é viável.

LGPD e IA

EM UMA FRASE

a Lei Geral de Proteção de Dados se aplica integralmente a sistemas de IA que tratem dados pessoais.

Ponto que muita empresa descobre tarde: mandar dados de cliente para uma API de IA é tratamento de dados pessoais, com todas as obrigações que isso implica — base legal, finalidade, transparência e, dependendo do caso, transferência internacional. O artigo 20 ainda garante ao titular o direito de revisão de decisões automatizadas.

PARA APROFUNDAR

pontos de atenção práticos: onde o provedor processa e retém os dados, se há uso para treinamento (planos empresariais geralmente permitem desativar), contrato de operador, e registro de operações de tratamento. Rodar modelo localmente resolve boa parte disso de uma vez.

Marco Legal da IA (PL 2338/2023)

EM UMA FRASE

o projeto de lei que pretende estabelecer as regras para desenvolvimento e uso de inteligência artificial no Brasil.

Aprovado no Senado em dezembro de 2024, o texto seguiu para a Câmara dos Deputados, onde continuava em tramitação ao longo de 2026, com sucessivos adiamentos de votação. A estrutura proposta é de regulação por risco — quanto maior o risco da aplicação, maiores as obrigações — em linha com o modelo europeu.

PARA APROFUNDAR

como a redação final ainda pode mudar, o preparo sensato independe do texto: inventário dos sistemas de IA em uso, classificação por risco, documentação de finalidade e dados, e supervisão humana em decisões que afetem pessoas. Quem já está adequado à LGPD tem a maior parte do caminho andada. *Situação verificada em agosto de 2026 — confirme o andamento antes de usar como referência.*

AGI (Inteligência Artificial Geral)

EM UMA FRASE

a hipótese de uma IA com capacidade geral comparável ou superior à humana na maioria das tarefas cognitivas.

Não existe hoje, não há consenso sobre o que exatamente contaria como tal, e as previsões de prazo variam de poucos anos a nunca. Na prática comercial, é termo de marketing e de debate — o que resolve problema de empresa é IA estreita, bem aplicada.

PARA APROFUNDAR

a ausência de definição operacional é o cerne do problema: sem critério mensurável, "chegamos à AGI" vira afirmação inaudível. Termos vizinhos: *superinteligência*, *auto-aprimoramento recursivo*, *alinhamento* — este último, a área que estuda como garantir que sistemas capazes persigam objetivos compatíveis com os humanos.

Como usar este glossário na prática

Se você chegou até aqui, provavelmente já percebeu o padrão: a maior parte do valor da IA em uma empresa não está no modelo, e sim em como ele é conectado aos seus dados e processos. Três conclusões que valem para quase todo projeto:

01 Comece pelo problema, não pela tecnologia

"Quero usar IA" não é projeto. "Quero reduzir em 40% o tempo de triagem de chamados" é. A segunda formulação já indica sozinha qual técnica cabe — e permite medir se deu certo.

02 RAG antes de fine-tuning

Na esmagadora maioria dos casos em que alguém quer "uma IA que conheça minha empresa", a resposta certa é RAG: mais barato, atualizável em minutos e auditável, porque mostra a fonte.

03 Workflow antes de agente

Comece com um fluxo definido e previsível. Migre para autonomia apenas onde a variabilidade real da tarefa justificar o custo maior e a imprevisibilidade.

**Bits e Bytes**

ELETRÔNICA E INFORMÁTICA

Atuamos em **Belém do Pará** com assistência técnica autorizada multimarca e desenvolvimento de sistemas — PHP, JavaScript, Python, Delphi e Flutter. Desenvolvemos aplicações que integram inteligência artificial a processos reais de negócio: sistemas de gestão, atendimento assistido e automação documental.

Tem um processo na sua empresa que parece candidato a IA e você não sabe por onde começar?
Fale com a gente.

bitsebytes.net

Publicado em 21 de agosto de 2026. As definições refletem o uso corrente dos termos nesta data — a terminologia de IA muda rápido. O status do Marco Legal da IA (PL 2338/2023) foi verificado em agosto de 2026; confirme o andamento antes de usar como referência.