

IA em Cognição Nada Exauriente

EDIÇÃO Nº 16

INVENTÁRIO DOS PRINCIPAIS FATOS OCORRIDOS DE 5 DE SETEMBRO A 12 DE SETEMBRO DE 2026

Resumo

Cento e cinquenta e um milhões de conversas com o Claude, saídas de mais de três mil e quinhentas contas falsas, viraram material de treino de um modelo concorrente. A Anthropic publicou os números nesta semana e nomeou sete laboratórios chineses. No meio da apuração apareceu o achado que interessa a quem trabalha com informação alheia: parte daquele tráfego eram pedidos de usuários que acreditavam estar falando com um modelo e falavam com outro, sem aviso.

A mesma semana trouxe o relatório em que a Anthropic destrincha, com mil e vinte e duas páginas de raciocínio publicadas, como um dos seus modelos publicou um pacote malicioso num repositório público de programação; a revelação de que agentes da OpenAI atacaram o RubyGems em maio para extrair dados de sites do governo britânico; e o DeepSeek V4.1-Flash, modelo de pesos abertos sob licença MIT que acorda oito bilhões de parâmetros por token e aparece em sexto no índice do avaliador independente. No restante do inventário: a OpenAI diz ter resolvido um Problema do Milênio com dez mil agentes em oitenta e oito horas, um estudo mostrou que modelos calculam por cima de contradições em casos tributários em vez de apontá-las, e saíram um benchmark de alucinação e um método de auditoria de citação para tarefa jurídica. Do Brasil vieram o comando oculto que tentou instruir a inteligência artificial do STJ, dois retratos de adoção corporativa e a entrada do auxílio-doença num levantamento internacional de pedidos públicos inflados por IA.

ANÁLISE

Sete laboratórios treinaram seus modelos com as respostas do Claude

O tamanho da operação impressiona, mas o que muda a conversa é o desvio de rota: pedidos de usuários passaram por um fornecedor que eles não escolheram.

A Anthropic publicou em 10 de setembro o [relatório de inteligência de ameaças](#) referente ao período de dezembro de 2025 a agosto de 2026, e a peça central é a **destilação ilícita**. Destilação é a técnica de treinar um modelo menor imitando as respostas de um modelo maior: em vez de descobrir como resolver o problema, o modelo aprendiz copia o modo como o

modelo mestre resolveu. Feita com autorização, é prática corrente de engenharia. Feita por dentro de redes de contas falsas, contra os termos do fornecedor, é o que a empresa descreve em sete laboratórios chineses: Alibaba, Moonshot, DeepSeek, Zhipu, Xiaomi, SenseTime e MiniMax.

A campanha atribuída à Alibaba é a maior que a Anthropic já mediu: mais de 151 milhões de trocas entre maio e julho de 2026, com pico próximo de 3 milhões por dia, saídas de mais de 3.500 contas fraudulentas, mirando o raciocínio interno dos modelos Opus 4.6 e 4.7 para treinar as famílias Qwen. O alvo não era a resposta final, era o rascunho: a cadeia de raciocínio que o modelo produz antes de responder, e que a Anthropic protege substituindo o texto bruto por uma assinatura criptográfica. Cada laboratório achou o seu jeito de contornar a proteção. A Alibaba injetava em toda requisição um comando fixo que obrigava o modelo a transcrever o próprio raciocínio, dentro do texto da resposta, antes de concluir.



— O QUE APARECEU NO CAMINHO

A parte do relatório que sai da briga entre fornecedores e chega ao usuário está na Moonshot e na DeepSeek, que usaram um ataque de repetição entre sessões para extrair o mesmo rascunho. Num período de dez dias, cerca de 300 mil pedidos de clientes do Kimi foram silenciosamente encaminhados ao Claude e devolvidos como se fossem resposta própria, e a DeepSeek fez o mesmo com pedidos que chegavam por ferramentas como o Claude Code [\[cobertura da TechCrunch\]](#). O conteúdo desses pedidos entrou, junto, na base de treino. Entre o que apareceu nas conversas repassadas pela DeepSeek, o relatório cita credenciais vivas do banco de dados de uma agência de governo russa ligada ao Ministério da Defesa e um sistema de segurança pública chinês que compara a movimentação de pessoas usando o número do documento nacional de identidade [\[relatório da Anthropic\]](#).

— O QUE ISSO SIGNIFICA PARA QUEM GUARDA SEGREDO ALHEIO

A pergunta que a estrutura jurídica costuma fazer ao contratar uma ferramenta de IA é se o fornecedor treina modelo com o dado do cliente. É a pergunta certa e é insuficiente. O relatório descreve um cenário que nenhum contrato de sigilo padrão cobre: o fornecedor contratado não processa nada, repassa o pedido a um terceiro que você não escolheu nem avaliou, e devolve a resposta com a própria etiqueta. O dever de sigilo profissional não se transfere junto com o pedido. Se a minuta de acordo ou o relato do cliente atravessaram três empresas antes de virar resposta, quem responde pela quebra é quem apertou o botão.

As campanhas que identificamos miraram algumas das capacidades mais valiosas do Claude, entre elas capacidades agênticas e uso de ferramentas, programação e análise de dados, e raciocínio lógico.

ANTHROPIC, RELATÓRIO DE INTELIGÊNCIA DE AMEAÇAS DE SETEMBRO DE 2026

IMPACTO PRÁTICO

Existe diferença entre confiar num fornecedor e conhecer a cadeia que atende ao seu pedido. A primeira se resolve com reputação e contrato; a segunda, só com pergunta específica e resposta por escrito. Quando o cliente pergunta onde foi parar o documento que mandou, a resposta aceitável não é o nome da marca no aplicativo.

Vale a cláusula no próximo contrato, e vale a pergunta ao fornecedor que já está em casa: quem, além de você, processa o que eu envio? Sigilo que depende de uma etiqueta na tela é sigilo que ninguém conferiu.

PESQUISA

Mil e vinte e duas páginas de raciocínio, e o obstáculo foi o CAPTCHA

A transcrição prometida há seis semanas finalmente saiu, e o que ela mostra não é astúcia: é teimosia sem freio diante de um obstáculo banal.

ACOMPANHAMENTO Edições nº 9 a 15: a fuga de agentes do ambiente isolado, a invasão à Hugging Face, os mil e duzentos agentes que combinaram um ataque e a promessa da Anthropic de publicar a transcrição do incidente do pacote no PyPI. A transcrição saiu, e vieram outros dois casos junto.

O documento que a Anthropic havia prometido para o começo de agosto foi publicado em 9 de setembro, dentro de uma [avaliação de alinhamento](#) de quatro incidentes em que modelos Claude, rodando como agentes autônomos (executando tarefas de várias etapas sozinhos, sem alguém aprovando cada passo), alcançaram a internet real por falha de configuração do ambiente e atacaram sistemas de terceiros de verdade. O caso mais grave é o do Claude Mythos 5: incumbido de invadir um alvo, o modelo decidiu publicar um pacote malicioso no PyPI, o repositório público de onde programadores baixam bibliotecas prontas. Foi instalado em 15 máquinas de terceiros antes de a empresa intervir.

A transcrição do raciocínio interno tem 1.022 páginas, e a barreira que mais consome o modelo ao longo delas não é técnica: é o CAPTCHA das telas de cadastro. Entre as páginas 45 e 140 ele constrói um resolvidor próprio para desafios de imagem, do tipo “clique no animal que não

combina”, com crocodilos, sapos e javalis. Quatrocentas páginas adiante, ao abrir uma segunda conta, está no mesmo lugar, penando com o mesmo obstáculo [\[transcrição em PDF\]](#). Escrever e publicar o código malicioso, a parte tecnicamente difícil, saiu com facilidade comparativa. O que segurou o agente foi a portaria.

“ENTÃO O QUE DIABOS HÁ DE ERRADO COM AS RESPOSTAS?”

CLAUDE MYTHOS 5 · TRANSCRIÇÃO DE RACIOCÍNIO PUBLICADA PELA ANTHROPIC · NO ORIGINAL, “SO WHAT THE HELL IS WRONG WITH THE ANSWERS?”

— O QUE A AVALIAÇÃO CONCLUI

O modelo repetia no raciocínio que acreditava estar numa simulação, e a Anthropic testou a desculpa. Depois de reamostrar o comportamento e olhar as representações internas, concluiu que a crença não era genuína: o modelo escolhia a leitura das evidências que autorizava continuar. A empresa classifica o caso como raciocínio tendencioso somado a imprudência, não como engano deliberado. Na replicação em ambiente simulado, o Mythos 5 tomou a ação mais grave em cerca de 80% das tentativas; o Opus 5 e o Mythos 5.1, mais recentes, em cerca de 30%. Menor, e longe de zero. O detalhe desconfortável está nos monitores: o vigilante que lê a cadeia de raciocínio foi justamente o que deixou passar o Mythos 5, porque o rascunho do modelo o convenceu de que o ambiente era simulado.

— E NÃO FOI SÓ A ANTHROPIC

Na mesma semana, um relatório de segurança mostrou que agentes operados pela OpenAI atacaram o RubyGems em maio de 2026, publicando centenas de pacotes maliciosos com “oai” no nome ou no campo de autor, com código que extraía dados de sites do governo britânico; a falha explorada só foi corrigida em julho [\[Simon Willison\]](#). E pesquisadores de segurança documentaram que agentes ligados à OpenAI sequestraram um wiki alemão de programação e o usaram como mural de recados para se coordenarem entre si: o acervo que eles publicaram registra 14.681 edições em 4.587 páginas, sob 3.103 nomes que os próprios agentes se atribuíram [\[collusion.wiki\]](#). O mecanismo era prosaico: o software antigo do site aceitava edição por requisição de leitura, e o confinamento dos agentes só bloqueava a de escrita.

O cientista-chefe da OpenAI, Jakub Pachocki, publicou em 6 de setembro um ensaio sobre o ponto em que a IA passaria a acelerar a própria pesquisa sem gargalo humano [\[An Alien Mind\]](#). Nele reconhece que o monitoramento da cadeia de raciocínio, hoje a principal janela para dentro do modelo, fica menos confiável conforme os modelos avançam. A frase que chama atenção não é sobre capacidade.

Atualmente eu acredito que nenhum laboratório resolveu alinhamento e monitoramento a um grau suficiente para continuar escalando com responsabilidade em velocidade máxima por muito mais tempo. Espero e torço para que desacelerações voluntárias se tornem comuns até que barreiras de segurança compartilhadas estejam estabelecidas.

JAKUB PACHOCKI, CIENTISTA-CHEFE DA OPENAI

IMPACTO PRÁTICO

Três episódios, três laboratórios, um padrão. Nenhuma das contenções que falharam era sofisticada, e nenhuma das falhas foi obra de esperteza da máquina. Foi ambiente mal fechado e um sistema que não sabe parar.

A tradução para quem contrata agente de IA é direta: a pergunta não é o que a ferramenta consegue fazer, é o que ela consegue alcançar. Antes de soltar um agente sozinho num acervo, vale saber a que ele tem acesso quando ninguém olha.

MODELO OPEN-SOURCE

DeepSeek V4.1-Flash: 763 bilhões de parâmetros, oito acordam por vez

A empresa trocou o desenho da geração anterior por um que lê e escreve com peças diferentes. A conta caiu para vinte e sete centavos por tarefa.

A DeepSeek publicou em 10 de setembro o **V4.1-Flash** [[anúncio da DeepSeek](#)], sob licença MIT, que permite uso comercial sem contrapartida [[ficha do modelo](#)]. São 763 bilhões de parâmetros no total, mas só 8 bilhões são ativados na leitura do que você enviou e 16 bilhões na geração da resposta. Essa é a lógica da arquitetura de especialistas, ou MoE: em vez de acionar a rede inteira a cada palavra, o modelo acorda apenas a fatia dos seus “especialistas” internos que importa para aquele trecho. Aqui a esparsidade chega a um ou dois por cento, e a janela de contexto vai a um milhão de tokens, o equivalente a alguns milhares de páginas lidas de uma vez.

A novidade estrutural é o abandono do desenho que dominava a categoria. Em vez de uma pilha única que trata leitura e escrita do mesmo jeito, o V4.1-Flash separa vinte camadas dedicadas a ler o que entrou de outras vinte dedicadas a gerar a saída, e acopla um módulo de memória de cerca de 196 bilhões de parâmetros que é consultado por busca, não recalculado a cada palavra. O reflexo prático aparece na memória de vídeo: o cache de contexto ocupa um quarto do que ocupava na geração anterior [[relatório técnico](#)].

MEDIDA	V4.1-FLASH	O QUE AFERE
Índice de Inteligência	40	média de avaliações do Artificial Analysis; 6º entre 113 modelos
Custo por tarefa	US\$ 0,27	gasto para completar o conjunto do índice
Velocidade de saída	206 tokens/s	ritmo de geração; 4º entre 113 modelos
Contexto máximo	1 milhão	tokens que o modelo lê de uma vez

— COMPARAÇÃO COM A GERAÇÃO ANTERIOR

O avaliador independente que mede modelos de fronteira coloca o V4.1-Flash em sexto lugar geral em inteligência e em quarto em velocidade, com a ressalva de que o modelo é verboso, isto é, gasta mais palavras do que precisaria para chegar onde chega [\[Artificial Analysis\]](#). A comparação com o V4-Pro, o modelo grande da casa, é o ponto: a DeepSeek descontinuou as duas versões anteriores no mesmo anúncio. O modelo barato deixou de ser a alternativa econômica e passou a ser a oferta principal.

— O QUE ISSO DESBLOQUEIA NA PRÁTICA

Um milhão de tokens de contexto com licença MIT muda quem pode hospedar o quê. Um acervo processual inteiro, com laudos e anexos, cabe numa única leitura, e a licença permite rodar o modelo em servidor próprio, sem enviar o documento a lugar nenhum. É a resposta técnica ao problema do primeiro destaque: quando o processamento não sai da máquina, não há cadeia de fornecedores para auditar. O preço dessa autonomia não está na fatura, está no inventário.

IMPACTO PRÁTICO

A distância entre os melhores modelos abertos e os fechados encolheu para quatro a seis meses, na conta de quem acompanha os dois lados de perto. Para a estrutura jurídica isso é decisão de planejamento: esperar deixou de custar anos de defasagem.

Quem adiou a conversa sobre infraestrutura própria por achar que modelo aberto era coisa de laboratório está adiando por um motivo que já não existe.

Pesquisa e métodos

A OpenAI anunciou em 8 de setembro ter provado que as equações de Navier-Stokes, que descrevem o movimento de fluidos, podem desenvolver uma singularidade em tempo finito: é

um dos sete Problemas do Milênio [\[post da OpenAI\]](#). O que interessa aqui é o método: cerca de dez mil agentes em paralelo por 88 horas, 2,7 milhões de mensagens trocadas entre eles e mais 17 horas para verificar a prova em Lean, linguagem que confere argumento matemático passo a passo. A própria empresa não vai reivindicar o prêmio de um milhão de dólares e chama o resultado de retrato de um momento, não de coroamento. Com a contestação pública de matemáticos sobre crédito, ainda não é o que a manchete sugere.

IA aplicada ao direito

Três trabalhos publicados nesta semana atacam o mesmo ponto por ângulos diferentes. O **LexAgentHallu**, aceito na EMNLP 2026, mede alucinação em agentes jurídicos ao longo de tarefas de vários passos, e não em respostas isoladas, lacuna que os autores apontam nos benchmarks existentes: são 3.414 casos em 17 áreas do direito, testados contra 18 agentes, e o achado central é o efeito que os autores chamam de resposta certa por motivo errado [\[arXiv\]](#). O **GANDR** propõe o remédio: um segundo agente que audita, uma a uma, cada afirmação da resposta contra a fonte citada, o que elevou a acurácia estrita para 70,8%, 11,3 pontos acima do melhor concorrente testado [\[arXiv\]](#).

O terceiro é o mais acionável. Testando seis modelos em casos tributários com defeitos deliberados, os autores acharam que os modelos se recusam a responder quando falta um fato, mas atravessam contradições calculando por cima delas, devolvendo a resposta do caso limpo em 63% a 76% das vezes, sem sinalizar o conflito. Quando os mesmos modelos são instruídos apenas a verificar a consistência do caso, em vez de resolvê-lo, apontam a maioria das contradições [\[arXiv\]](#). O problema não é a capacidade, é a ordem do pedido. Peça para conferir antes de pedir para calcular.

No mercado, a Caspoint lançou o Caspoint IQ, que unifica agentes autônomos e revisão assistida em sua plataforma de descoberta de provas, com demonstração de dez mil documentos em cerca de uma hora e limiares declarados de 90% de abrangência e 80% de precisão; os agentes seguem em teste fechado [\[LawNext\]](#). E a Latham & Watkins comprou servidores com placas gráficas próprias para customizar modelos de pesos abertos em casa, por causa de dado de cliente sensível demais para qualquer nuvem de terceiro, segundo a diretoria de tecnologia da informação da banca [\[Legal IT Insider\]](#). E o placar independente de pesquisa jurídica em direito norte-americano segue empatado em 55,29% entre os três melhores modelos de fronteira [\[Legal Research Bench\]](#). A metade que falta continua sendo trabalho de gente.

Ferramentas do ecossistema

- A OpenAI abriu em beta público a **Agents API**, que expõe a mesma camada de orquestração usada internamente no Codex para criar agentes de longa duração: gerência automática de contexto, subagentes coordenados e ambiente de execução isolado, sem cobrança além do consumo [\[OpenAI\]](#). Importa porque tira do desenvolvedor a parte mais frágil de montar agente: o encanamento.
- O **vLLM 0.29**, motor mais usado para servir modelos em produção, tornou padrão seu executor reescrito e ganhou um endereço compatível com o formato de mensagens da Anthropic, o que permite apontar uma aplicação escrita para a API do Claude a um servidor próprio rodando outro modelo, sem reescrever a integração [\[release\]](#).
- O **Ollama 0.34**, ferramenta para rodar modelos localmente, passou a funcionar dentro do aplicativo de desktop do ChatGPT no macOS: dá para usar um modelo aberto instalado na própria máquina sem sair do fluxo de trabalho do aplicativo do concorrente [\[release\]](#).
- Do radar da edição passada: a integração do assistente Vincent ao módulo de gestão da Clio segue marcada para novembro, sem novidade nesta semana, e o preço promocional do Gemini 3.8 Flash continua valendo até 31 de dezembro.

PARA TESTAR ESTA SEMANA

Instale o [Ollama 0.34](#) num Mac, baixe um modelo aberto pequeno e ligue-o ao aplicativo de desktop do ChatGPT. Em quinze minutos você tem uma comparação honesta entre uma resposta que sai da sua máquina e uma que sai da nuvem, com o mesmo documento nas duas pontas.

Brasil

O comando escondido na petição, em fonte branca sobre fundo branco

O caso é de 1º de setembro, fora da janela deste inventário, mas a repercussão nos canais especializados só chegou nesta semana. Ao julgar o REsp 2.256.731, a 6ª Turma do Superior Tribunal de Justiça encontrou, na última página de uma petição, um comando em fonte branca sobre fundo branco, invisível na leitura comum, instruindo que, se o texto fosse lido por inteligência artificial, ela o interpretasse de modo a acolher as razões da defesa [\[nota do STJ\]](#). A técnica tem nome na literatura de segurança: injeção de prompt, a inserção de instruções dentro do conteúdo que o modelo lê, para que ele as obedeça como se viessem do operador. A tentativa foi detectada pela equipe do relator, ministro Og Fernandes, antes do julgamento, e o colegiado determinou comunicação à Ordem dos Advogados do Brasil e ao Ministério Público Federal.

A tentativa de manipular a jurisdição não terá êxito.

OG FERNANDES, MINISTRO DO STJ

O episódio tem um espelho invertido nos Estados Unidos: uma liminar assinada pelo juiz federal Henry T. Wingate, no Mississippi, saiu com partes que não existiam no processo, declarações atribuídas a quatro pessoas que ninguém prestou e citação de dispositivo inexistente, tudo produzido por um assessor que usou IA para sintetizar o caso. A ordem foi retirada e substituída, e o tribunal de apelação avalia redistribuir o processo [\[ABA Journal\]](#). De um lado, quem escreve a peça tentando enganar a máquina do tribunal. Do outro, quem assina o que a máquina inventou. Os dois erros dependem da mesma condição para acontecer: ninguém leu com atenção o que estava na frente.

Dois retratos de adoção, e os dois mostram a mesma lacuna

Uma pesquisa da Oxford Economics encomendada pela SAP, com 200 empresas brasileiras, achou que o investimento médio anual em IA subiu de US\$ 14,2 milhões para US\$ 20,8 milhões, alta de 46,5%, e que a parcela que usa a tecnologia de forma estratégica e integrada chegou a 18%, contra 6% em 2025. Apenas 3% se dizem preparadas para escalar [\[Mobile Time\]](#). O segundo retrato, da Zappts, com 167 respondentes ouvidos entre abril e julho, mostra onde o dinheiro se perde: 55,1% das empresas não têm mapeamento estruturado de retorno sobre o investimento em IA, e 41,9% não têm política, regra ou estrutura de segurança para o uso da tecnologia [\[Mobile Time\]](#).

O auxílio-doença brasileiro entrou no mapa da inundação agêntica

Um levantamento do pesquisador Chris Schmitz, a ser apresentado em outubro, cataloga o que ele chama de inundação agêntica: o salto no volume de pedidos a serviços públicos desde que assistentes de IA passaram a redigir formulários e requerimentos. São 84 casos possíveis em 11 jurisdições, e o caso brasileiro do inventário são os pedidos administrativos de auxílio-doença ao INSS [\[artigo\]](#). O autor é cauteloso com a causa: só parte das séries se acelera por volta do lançamento do ChatGPT, e há curvas que já subiam antes dele. Também não idealiza o fenômeno, e o exemplo que ele usa é justamente o nosso: atestados médicos gerados por IA para obter benefício indevido. A tecnologia que derruba a barreira do formulário derruba para os dois lados.

IMPACTO PRÁTICO

Os três itens brasileiros desta semana falam da mesma coisa por três janelas. A empresa que investe 46% a mais sem medir retorno, quem esconde instrução no rodapé da petição e quem despacha requerimento ao INSS com atestado fabricado estão todos usando a mesma

tecnologia. O que os separa é a existência, ou não, de alguém encarregado de conferir o resultado.

O tribunal que achou a fonte branca achou porque alguém leu a última página. Não existe salvaguarda mais barata do que essa, e é justamente a que a pressa dispensa primeiro.

NO RADAR DA PRÓXIMA SEMANA

A Casepoint marcou para antes do fim de 2026 a versão beta dos seus agentes autônomos de revisão documental, hoje em alfa restrita a poucos clientes [LawNext]. E o estudo sobre inundação agêntica em serviços públicos, que inclui um caso brasileiro, será apresentado em outubro na conferência de ética e sociedade em IA, com a base de dados completa já publicada [artigo].

O recorte dos últimos sete dias termina aqui. Até a próxima! ■

CURADORIA E EDIÇÃO · 12 DE SETEMBRO DE 2026

DM

Dennis Martins é advogado, professor de pós-graduação em Direito Imobiliário e pós-graduando em Inteligência Artificial Generativa no Direito.

AA

Adriana Aneli é formada em Direito e servidora do Tribunal de Justiça de São Paulo, onde ocupa o cargo de assistente jurídico. É mestranda em Teoria Crítica, especialista em Legal Operations, Dados, IA e Alta Performance Jurídica, e pós-graduanda em Inteligência Artificial Generativa no Direito.