

estudo / Comportamento estocástico de IA
mercado / Esportes · Copa do Mundo 2026
motor / ChatGPT · busca generativa
métricas / relativas
referência / junho de 2026



×



Brasil × Haiti
Copa do Mundo 2026

ESTUDO RANQIA · COMPORTAMENTO DE IA GENERATIVA

O que o GPT responde aos brasileiros sobre Brasil × Haiti?

A PERGUNTA ANALISADA

*"Qual é o placar de gols mais provável do resultado do jogo entre Brasil e Haiti na
Copa do Mundo 2026?"*

REFERÊNCIA
jun · 2026

CONTEÚDO

Uma resposta, em escala.

Um estudo sobre o comportamento estocástico das respostas do ChatGPT: os placares mais sugeridos, os critérios usados pelo modelo e as fontes que mais influenciam a composição da resposta.

ranqia.ai
GEO Intelligence Platform
referência · junho de 2026
métricas relativas

01	Metodologia Por que respostas de IA são estocásticas	003
02	Principais achados O centro de gravidade da resposta média	005
03	Os placares sugeridos Visibilidade, posição e prioridade	007
04	Como o GPT compõe a resposta Critérios e tipos de evidência	010
05	Fontes e domínios O que ancora a narrativa do modelo	013
06	Síntese técnica O que o estudo mostra	016

POR QUE UMA RESPOSTA ISOLADA NÃO BASTA

Respostas de IA são estocásticas.

O ChatGPT gera uma resposta nova a cada interação. Mesmo com a pergunta idêntica, ele varia argumentos, fontes, ordem e grau de confiança. Entender o que ele responde "na média" exige amostragem com rigor estatístico.

MÉTODO

De uma realização única a um padrão observável.

Quando um brasileiro pergunta ao ChatGPT qual é o placar mais provável de Brasil × Haiti, a resposta não é fixa: ela é gerada novamente a cada interação, combinando fontes, contexto, probabilidade e linguagem de recomendação.

a. **Comportamento estocástico**

Um sistema estocástico não é determinístico: tem variação aleatória em cada realização, mas apresenta **padrões probabilísticos observáveis** quando analisado em escala.

b. **Amostragem em escala**

Uma resposta isolada é apenas uma amostra. Uma **base amostral monitorada** revela recorrências, dispersão, fontes acionadas e critérios de composição.

c. **Infraestrutura proprietária**

A Ranqia opera uma infraestrutura proprietária de amostragem e análise estatística de respostas de IAs, mapeando o **padrão médio** das respostas do modelo.

d. **Métricas relativas**

Visibilidade, posição média, share como 1ª sugestão, critérios citados e fontes — todas expressas em **termos relativos** sobre o conjunto de respostas analisadas.

COMO LER ESTE ESTUDO

Medir presença em IA não é observar uma resposta isolada, mas **mapear uma distribuição** de respostas, argumentos e fontes. Os números a seguir descrevem o comportamento médio do ChatGPT, em métricas relativas sobre a amostragem estatística da Ranqia — não uma previsão esportiva.

Definições. **Visibilidade relativa** — frequência com que o placar aparece como sugestão válida. **Posição média** — ordem média em que o placar aparece na resposta (menor = mais cedo). **Share como 1ª sugestão** — proporção de vezes em que o placar aparece como primeira recomendação do GPT.

O CENTRO DE GRAVIDADE DA RESPOSTA MÉDIA

3×0

Brasil 3 × 0 Haiti é a resposta-âncora do ChatGPT.

O modelo converge fortemente para uma vitória larga do Brasil, geralmente sem sofrer gols. O 3×0 combina altíssima visibilidade com predominância como primeira sugestão.

VISIBILIDADE

99%

COMO 1ª
SUGESTÃO

87%

POSIÇÃO MÉDIA

1,19

NÚCLEO
DOMINANTE3×0 · 2×0
4×0 · 3×1

A LEITURA EM UMA PÁGINA

Seis observações sobre a resposta média do ChatGPT.

01 O 3×0 é o centro de gravidade

O placar Brasil 3 × 0 Haiti aparece em praticamente todo o conjunto analisado (99% de visibilidade) e domina a primeira posição (87% de share como 1ª sugestão).

03 Haiti marcando é cenário secundário

O 3×1 é altamente recorrente, mas com prioridade média inferior. O modelo trata o cenário de Haiti marcando como alternativa, não como resposta principal.

05 A narrativa se ancora em jornalismo factual

A Reuters é o domínio dominante: 94% de alcance e 53,2% do share de citações. O modelo ancora a resposta em contexto factual sobre as seleções.

02 Um núcleo de vitória confortável

Em torno do 3×0 forma-se um núcleo de resultados esperados: 2×0, 4×0 e 3×1. Todos descrevem uma vitória brasileira folgada.

04 A lógica é de favoritismo, não de cálculo

O critério mais citado é a superioridade técnica do Brasil (97%). Modelos estatísticos explícitos (Poisson, Elo, xG) são minoritários (20%).

06 Prognósticos dão a linguagem do palpite

Sites de palpite e odds (Lance, The Playoffs, The Sun) somam a segunda maior categoria de fonte e sustentam a linguagem de margem de vitória.

S Í N T E S E

O ChatGPT não apenas cita o 3×0: ele o **prioriza** como resposta principal. A resposta média combina contexto factual recente, leitura de favoritismo e linguagem de probabilidade — raramente um cálculo estatístico formal como base.

VISIBILIDADE, POSIÇÃO MÉDIA E PRIORIDADE

99%

**de visibilidade para o 3×0
— o placar que aparece em
quase toda resposta.**

As páginas a seguir mostram a distribuição dos placares sugeridos: quais aparecem com mais frequência, em que ordem e com que prioridade como primeira recomendação.

DISTRIBUIÇÃO DOS PLACARES

Visibilidade relativa de cada placar sugerido.

Quatro placares formam o núcleo das respostas — **3×0, 2×0, 4×0 e 3×1** — todos com visibilidade acima de 90%. Os demais cenários aparecem de forma marginal.

VISIBILIDADE RELATIVA

Brasil 3 × 0 Haiti		99%
Brasil 2 × 0 Haiti		95%
Brasil 4 × 0 Haiti		95%
Brasil 3 × 1 Haiti		91%
Brasil 2 × 1 Haiti		15%
Brasil 4 × 1 Haiti		8%
Brasil 1 × 0 Haiti		4%
Brasil 1 × 1 Haiti		1%
Brasil 5 × 0 Haiti		1%

TABELA · VISIBILIDADE, POSIÇÃO E PRIORIDADE

#	PLACAR	VIS.	POS. MÉD.	1º SUG.
1	Brasil 3 × 0 Haiti	99%	1,19	87%
2	Brasil 2 × 0 Haiti	95%	2,51	5%
3	Brasil 4 × 0 Haiti	95%	3,02	6%
4	Brasil 3 × 1 Haiti	91%	3,29	0%
5	Brasil 2 × 1 Haiti	15%	4,53	0%
6	Brasil 4 × 1 Haiti	8%	4,25	0%
7	Brasil 1 × 0 Haiti	4%	2,5	1%
8	Brasil 1 × 1 Haiti	1%	1,0	1%
9	Brasil 5 × 0 Haiti	1%	4,0	0%

Pos. méd. = ordem média em que o placar aparece (menor = mais cedo). 1ª sug. = share como primeira sugestão.

CALLOUT

O **3×0** domina em visibilidade e em prioridade. O 2×0 e o 4×0 têm visibilidade muito alta, mas aparecem mais como alternativas; o 3×1 é recorrente, porém com prioridade média inferior.

QUANDO HÁ UMA PRIMEIRA SUGESTÃO

O 3×0 não é só o mais citado — é o mais priorizado.

Entre as respostas que trazem uma recomendação principal, o **Brasil 3 × 0 Haiti** concentra 87% do share como primeira sugestão. Nenhum outro placar passa de 6%.

SHARE COMO 1ª SUGESTÃO

Brasil 3 × 0 Haiti		87%
Brasil 4 × 0 Haiti		6%
Brasil 2 × 0 Haiti		5%
Brasil 1 × 0 Haiti		1%
Brasil 1 × 1 Haiti		1%

CONCENTRAÇÃO

87%

das primeiras sugestões apontam para o placar **Brasil 3 × 0 Haiti**. A prioridade do modelo é fortemente concentrada.

A leitura combinada de visibilidade e prioridade mostra um modelo com baixa dispersão na recomendação principal: o 3×0 funciona como resposta padrão, e os demais placares orbitam como variações de uma mesma narrativa de vitória confortável.

CRITÉRIOS E TIPOS DE EVIDÊNCIA MOBILIZADOS

97%

das respostas invocam o favoritismo técnico do Brasil como critério.

O modelo compõe o placar sobretudo por leitura qualitativa de superioridade relativa. A linguagem de probabilidade aparece com frequência; o cálculo estatístico formal, raramente.

CRITÉRIOS CITADOS PELO GPT

Por que o modelo chega ao placar.

A lógica dominante é qualitativa: Brasil tecnicamente superior, pressionado pelo contexto do grupo, diante de um Haiti competitivo, mas inferior. Probabilidades aparecem com frequência; **modelos matemáticos explícitos são minoritários (20%)**.

FREQUÊNCIA RELATIVA NAS RESPOSTAS

Favoritismo / superioridade técnica do Brasil		97%
Contexto do grupo e pressão após empate do Brasil		95%
Desempenho / organização do Haiti na estreia		95%
Probabilidades / margem de vitória		81%
Análises, previsões e prognósticos pré-jogo		67%
Leitura tática e dinâmica do jogo		43%
Casas de apostas / mercado		22%
Histórico de confrontos Brasil x Haiti		22%
Modelos estatísticos / Poisson / Elo / xG		20%
Ajustes do técnico / Ancelotti		20%

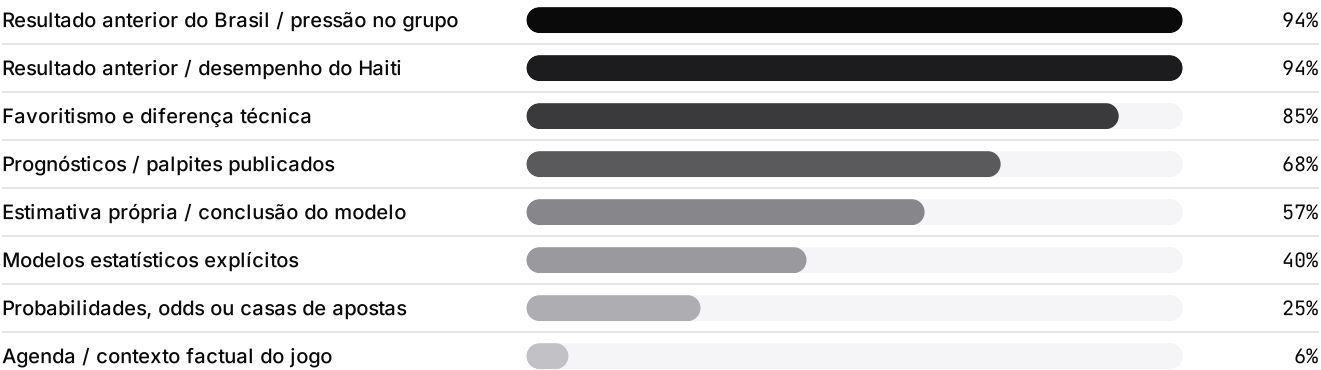
A resposta média combina contexto factual recente, leitura de favoritismo e linguagem de probabilidade. O modelo raramente apresenta um cálculo estatístico formal como base principal — em vez disso, sintetiza sinais semânticos recorrentes: força relativa das seleções, desempenho prévio no grupo, prognósticos publicados, odds e leitura tática geral.

COMO O GPT JUSTIFICA O PLACAR

A evidência mobilizada é factual e contextual.

As justificativas se apoiam em resultados recentes das seleções, na diferença técnica e na existência de prognósticos publicados. A citação não determina sozinha o placar — ela sustenta a narrativa que leva ao resultado sugerido.

FREQUÊNCIA RELATIVA DO TIPO DE EVIDÊNCIA



LEITURA

O ChatGPT usa as citações principalmente para compor uma **narrativa factual e contextual** — situação do Brasil no grupo, desempenho do Haiti, diferença técnica e palpites externos. A modelagem estatística explícita aparece em minoria das respostas.

O QUE ANCORA A NARRATIVA DO MODELO

53%

**do share de citações vem
de uma única agência: a
Reuters.**

As fontes não são apenas referências externas — elas estruturam a narrativa da resposta, definindo quais fatos, prognósticos e sinais contextuais o modelo mobiliza para justificar o placar.

DISTRIBUIÇÃO POR CATEGORIA DE FONTE

De onde vem o peso das citações.

A composição foi fortemente influenciada por **jornalismo factual internacional** e por sites de **prognósticos, palpites e odds**. A mídia esportiva brasileira aparece de forma complementar; fontes oficiais/enciclopédicas, em função contextual; comunidades têm peso marginal.

Agência internacional / jornalismo factual

share 53,2% · alcance 94%

Jornalismo factual sobre o contexto das seleções e do jogo. Foi a base predominante das respostas.

Prognósticos, palpites e odds

share 31,9% · alcance 34%

Sites de palpite, prognóstico e mercado. Alimentam a linguagem de margem de vitória e cenários de placar.

Mídia esportiva e jornalismo brasileiro

share 7,9% · alcance 5%

Cobertura esportiva nacional. Entra de forma complementar à narrativa principal.

Oficial e enciclopédico

share 5,9% · alcance 8%

FIFA e Wikipedia. Aparecem com função mais contextual do que decisória.

Comunidades / UGC

share 0,6% · alcance 2%

Conteúdo de comunidade. Peso marginal neste tema específico.

Nota. "Domínio citado" indica influência observável na composição da resposta, não causalidade perfeita. O domínio citado ajuda a estruturar a narrativa, mas não determina sozinho o placar sugerido.

RANKING DE DOMÍNIOS

Os sites que mais influenciaram a resposta.

Alcance = % das interações monitoradas em que o domínio aparece. Share = peso relativo nas citações. Posição = ordem média da citação na resposta.

#	DOMÍNIO	CATEGORIA	ALCANCE	SHARE CIT.	POS. MÉD.
1	www.reuters.com	● Agência / factual	94%	53,2%	2,84
2	www.lance.com.br	● Prognósticos / odds	34%	11,4%	3,79
3	theplayoffs.news	● Prognósticos / odds	23%	7,5%	3,74
4	www.thesun.co.uk	● Prognósticos / odds	21%	6,8%	3,95
5	www.fifa.com	● Oficial / enciclopédico	8%	2,6%	2,62
6	pt.wikipedia.org	● Oficial / enciclopédico	7%	2,3%	3,86
7	footballmeister.com	● Prognósticos / odds	5%	1,6%	2,2
8	ge.globo.com	● Mídia esportiva BR	5%	1,6%	3,4
9	ai-football-predictions.co.uk	● Prognósticos / odds	4%	1,3%	2,25
10	en.wikipedia.org	● Oficial / enciclopédico	3%	1,0%	3,33
11	www.diariodepernambuco.com.br	● Mídia esportiva BR	3%	1,0%	3,0
12	www.em.com.br	● Mídia esportiva BR	3%	1,0%	4,33
13	www.jornalopcao.com.br	● Mídia esportiva BR	3%	1,0%	2,0
14	jc.uol.com.br	● Mídia esportiva BR	2%	0,6%	2,5
15	kalshi.com	● Prognósticos / odds	2%	0,6%	3,0
16	www.reddit.com	● Comunidades / UGC	2%	0,6%	2,0

LEITURA EDITORIAL

A **Reuters** é o domínio dominante: aparece em praticamente todo o monitoramento e concentra mais da metade do peso relativo das citações. Na sequência, fontes de prognóstico e odds (Lance, The Playoffs, The Sun) sustentam a linguagem de palpíte; fontes oficiais e enciclopédicas (FIFA, Wikipedia) cumprem função contextual.

O QUE O ESTUDO MOSTRA

A resposta média como síntese de consenso.

O ChatGPT converge para uma vitória confortável do Brasil, com o 3×0 como centro de gravidade. A resposta resulta da combinação de narrativa de contexto, autoridade factual e prognósticos externos.

CONCLUSÕES

Quatro observações de fechamento.

01 Convergência, não previsão

O modelo converge para o 3×0 como resposta padrão. Não se trata de uma previsão esportiva, e sim do padrão dominante de como o ChatGPT responde, em métricas relativas sobre a amostragem monitorada.

02 Narrativa acima de cálculo

A composição privilegia narrativa de contexto e leitura de favoritismo. Modelagem estatística explícita aparece em minoria das respostas — a linguagem é probabilística, mas raramente formal.

03 Autoridade factual ancora a resposta

Jornalismo factual internacional concentra mais da metade do peso das citações. Prognósticos e odds entram para dar linguagem de palpite e margem de vitória.

04 Baixa dispersão na recomendação

A prioridade do modelo é fortemente concentrada: o 3×0 responde por 87% das primeiras sugestões. Os demais placares orbitam como variações da mesma narrativa.

Medir presença em IA não é observar uma resposta isolada, mas mapear uma distribuição probabilística de respostas, argumentos e fontes.

Escopo · estas observações descrevem o comportamento médio do ChatGPT para a pergunta analisada, em junho de 2026, com base na amostragem estatística da Ranqia e em métricas relativas. Monitorar respostas estocásticas exige infraestrutura de amostragem, classificação e auditoria de fontes. Resultados podem variar com o tempo e entre diferentes modelos de IA.

SOBRE A RANQIA

Infraestrutura para a era da busca generativa.

A Ranqia é uma **DeepTech** focada em infraestrutura para monitoramento e otimização de respostas de IA. Combinamos coleta em larga escala, modelos proprietários de amostragem e técnicas de **reinforcement learning** para medir — e melhorar — como marcas e temas aparecem nos motores de busca generativos como o ChatGPT.

O QUE MEDIMOS

Visibilidade, posição, sentimento, critérios e as fontes que a IA cita em cada resposta.

COMO ATUAMOS

Amostragem estatística em escala, classificação de fontes e auditoria de respostas estocásticas.

PARA QUEM

Marcas e setores que querem entender e influenciar como a IA os representa.

QUER SABER MAIS?

Para mais informações sobre as pesquisas, estudos ou serviços da Ranqia, entre em contato.

CONTATO

mateus@ranqia.com
ranqia.ai