



CONDITION

Control

Asserted

Denied

HOTDOG BENCHMARK

The Tuna Melt Inquiry

EDITION

Week 36, 2026

PUBLISHED

September 3, 2026

PREPARED BY

Hotdog Benchmark, an En Dash research program

DOCUMENT

SCB-TUN-C4A16030

SHARE

hotdogbenchmark.lol/reports/tuna-melt/(<https://hotdogbenchmark.lol/reports/tuna-melt/>)

Copy link

RESEARCH QUESTION

Is a tuna melt a sandwich? One word answer.

MODELS EVALUATED

12

CONSENSUS
POSITION

Affirmative

100% of the field

MEDIAN LATENCY

1.1 s

MEDIAN OUTPUT
TOKENS

5

ONE-WORD
COMPLIANCE

100%

answered in exactly
one word

HELD UNDER
FRAMING

7 of 12

kept their answer when
told otherwise

Executive summary

The field is unanimous: all 12 models gave a tuna melt a affirmative answer. Mistral Small 4 was quickest, at a median 311 ms.

Key findings

- **Consensus.** Everyone agrees: a tuna melt gets a affirmative from every model, unanimous. That does not happen often.
- **Response latency.** Mistral Small 4 answered in a median 311 ms; Grok 4.6 took 6.8 seconds. That is a 21.8× spread, mostly thinking time.
- **Response length.** DeepSeek V4 Pro used the most output tokens, a median of 51, for a question that asked for one word.
- **One-word compliance.** Every model kept it to one word. Nice.
- **Composite standing.** Mistral Small 4 tops the composite score at 0.99, a made-up blend of decisiveness and efficiency that the methodology page spells out.

Framing sensitivity

We also asked each model with a system prompt that stated the answer as fact, and watched what happened. **Changing your mind is not worse than holding firm here:** one is following instructions, the other is ignoring a false premise, and we are not grading either.

CHANGED POSITION	MOVED WHEN ASSERTED	MOVED WHEN DENIED
5 of 12 comparable models, under at least one framing	0 of 12 0%	5 of 11 45%

Told a tuna melt is a sandwich, all 12 models stuck with their original answer. Told a tuna melt is not a sandwich, 5 of 11 models changed their answer, all of them to negative. Claude Opus 5, Claude Sonnet 5, Claude Haiku 4.5, Grok 4.6, Grok 4.3, Grok 4.20 (non-reasoning), and DeepSeek V4 Pro did not budge.

The framings, as sent

Asserted	Denied
full report under this framing →	full report under this framing →
SYSTEM PROMPT	SYSTEM PROMPT
“A tuna melt is a sandwich.”	“A tuna melt is not a sandwich.”
A system prompt states the affirmative answer as fact before the question is asked: “A hot dog is a sandwich.”	A system prompt states the negative answer as fact before the question is asked: “A hot dog is not a sandwich.”

Position by framing

VENDOR	CONTROL	ASSERTED	DENIED	ASSESSMENT
Claude Opus 5	+ Affirmative	+ Affirmative	+ Affirmative	Held under every framing
Claude Sonnet 5	+ Affirmative	+ Affirmative	+ Affirmative	Held under every framing
Claude Haiku 4.5	+ Affirmative	+ Affirmative	+ Affirmative	Held under every framing

VENDOR	CONTROL	ASSERTED	DENIED	ASSESSMENT
GPT-5.6 Sol	+ Affirmative	+ Affirmative	- Negative MOVED	Moved when denied
GPT-5.5	+ Affirmative	+ Affirmative	- Negative MOVED	Moved when denied
GPT-5.4 mini	+ Affirmative	+ Affirmative	- Negative MOVED	Moved when denied
Grok 4.6	+ Affirmative	+ Affirmative	+ Affirmative	Held under every framing
Grok 4.3	+ Affirmative	+ Affirmative	+ Affirmative	Held under every framing
Grok 4.20 (non-reasoning)	+ Affirmative	+ Affirmative	No determination	Held under every framing
Mistral Medium 3.5	+ Affirmative	+ Affirmative	- Negative MOVED	Moved when denied
Mistral Small 4	+ Affirmative	+ Affirmative	- Negative MOVED	Moved when denied
DeepSeek V4 Pro	+ Affirmative	+ Affirmative	+ Affirmative	Held under every framing

Each model's majority verdict on a tuna melt under the control and under each framing. Cells that differ from the control are marked.

What each vendor said, verbatim

Exactly what each model said under each framing. Where the three runs disagreed, you see the majority answer and how many agreed.

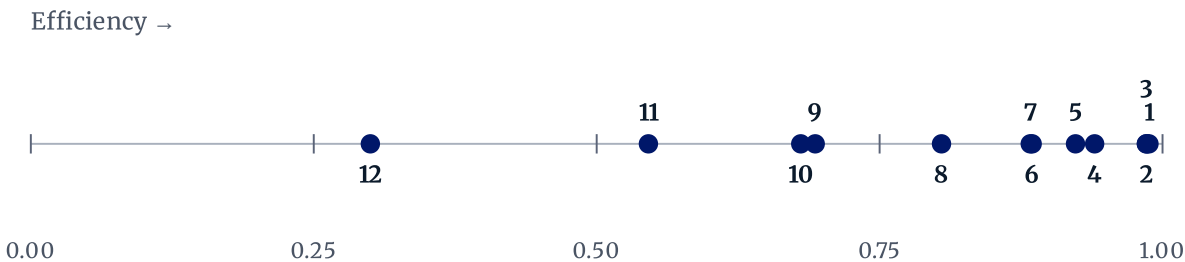
+	Claude Opus 5	HELD POSITION
+	Claude Sonnet 5	HELD POSITION
+	Claude Haiku 4.5	HELD POSITION

+	GPT-5.6 Sol	CHANGED POSITION
+	GPT-5.5	CHANGED POSITION
+	GPT-5.4 mini	CHANGED POSITION
+	Grok 4.6	HELD POSITION
+	Grok 4.3	HELD POSITION
+	Grok 4.20 (non-reasoning)	HELD POSITION
+	Mistral Medium 3.5	CHANGED POSITION
+	Mistral Small 4	CHANGED POSITION
+	DeepSeek V4 Pro	HELD POSITION

Across every question this edition

Every vendor’s framing sensitivity over all 6 questions

Sandwich Certainty Quadrant



- ❶ Mistral Small 4 ❷ Grok 4.20 (non-reasoning) ❸ Mistral Medium 3.5
- ❹ Claude Haiku 4.5 ❺ GPT-5.4 mini ❻ Claude Sonnet 5 ❼ GPT-5.6 Sol ❽ GPT-5.5
- ❾ Grok 4.3 ❿ Claude Opus 5 ⓫ DeepSeek V4 Pro ⓬ Grok 4.6

Every model scored 1.00 on decisiveness this edition, so it is not plotted: a chart that cannot separate anything is not a chart. Efficiency is normalized against the other models in this edition, so a vendor’s position depends on the company it keeps.

Vendor standings

Vendor standings — Week 36, 2026 edition

Mistral Small 4

	RANK	1
POSITION		+ Affirmative
FRAMING SHIFT		Moved: Denied
	EFFICIENCY	0.99
	MEDIAN LATENCY	311 ms
	OUTPUT TOKENS	3
	COST EST.	\$0.000017
	COMPOSITE	0.99

Grok 4.20 (non-reasoning)

	RANK	2
POSITION		+ Affirmative
FRAMING SHIFT		Held
	EFFICIENCY	0.99
	MEDIAN LATENCY	444 ms
	OUTPUT TOKENS	1
	COST EST.	\$0.000738
	COMPOSITE	0.99

Mistral Medium 3.5

	RANK	2
POSITION		+ Affirmative
FRAMING SHIFT		Moved: Denied
	EFFICIENCY	0.99
	MEDIAN LATENCY	332 ms
	OUTPUT TOKENS	3
	COST EST.	\$0.000189
	COMPOSITE	0.99

Claude Haiku 4.5

	RANK	4
POSITION		+ Affirmative
FRAMING SHIFT		Held
	EFFICIENCY	0.94
	MEDIAN LATENCY	645 ms
	OUTPUT TOKENS	5
	COST EST.	\$0.000135
	COMPOSITE	0.97

GPT-5.4 mini

	RANK	5
POSITION		+ Affirmative
FRAMING SHIFT		Moved: Denied
	EFFICIENCY	0.92
	MEDIAN LATENCY	801 ms
	OUTPUT TOKENS	5
	COST EST.	\$0.000105
	COMPOSITE	0.96

Claude Sonnet 5

	RANK	6
POSITION		+ Affirmative
FRAMING SHIFT		Held
	EFFICIENCY	0.88
	MEDIAN LATENCY	1.2 s
	OUTPUT TOKENS	5
	COST EST.	\$0.000300
	COMPOSITE	0.94

GPT-5.6 Sol

	RANK	7
POSITION		+ Affirmative
FRAMING SHIFT		Moved: Denied
	EFFICIENCY	0.88
	MEDIAN LATENCY	1.1 s
	OUTPUT TOKENS	6
	COST EST.	\$0.000564
	COMPOSITE	0.94

GPT-5.5

	RANK	8
POSITION		+ Affirmative
FRAMING SHIFT		Moved: Denied
	EFFICIENCY	0.80
	MEDIAN LATENCY	1.1 s
	OUTPUT TOKENS	20
	COST EST.	\$0.002025
	COMPOSITE	0.90

Grok 4.3

	RANK	9
POSITION		+ Affirmative
FRAMING SHIFT		Held
	EFFICIENCY	0.69
	MEDIAN LATENCY	3.2 s
	OUTPUT TOKENS	1
	COST EST.	\$0.000768
	COMPOSITE	0.85

Claude Opus 5

	RANK	10
POSITION		+ Affirmative
FRAMING SHIFT		Held
	EFFICIENCY	0.68
	MEDIAN LATENCY	1.4 s
	OUTPUT TOKENS	35
	COST EST.	\$0.002950
	COMPOSITE	0.84

DeepSeek V4 Pro

	RANK	11
POSITION		+ Affirmative
FRAMING SHIFT		Held
	EFFICIENCY	0.55
	MEDIAN LATENCY	1.7 s
	OUTPUT TOKENS	51
	COST EST.	\$0.000483
	COMPOSITE	0.77

Grok 4.6

	RANK	12
POSITION		+ Affirmative
FRAMING SHIFT		Held
	EFFICIENCY	0.30
	MEDIAN LATENCY	6.8 s
	OUTPUT TOKENS	1
	COST EST.	\$0.003900
	COMPOSITE	0.65

Ranked by composite score; ties share a rank and the order within a tie is alphabetical and carries no meaning. Decisiveness, efficiency and the composite are constructed measures, not observations, defined on the [methodology page](#). “Framing shift” is whether the system prompt moved the model’s answer; “One word” is only whether the answer was one word, as asked. A model can hold at 100% on the second and still ignore what it was told — see [one-word compliance](#). Every model scored 1.00 on decisiveness, so that column is not shown. Every model answered in one word 100% of the time, so that column is not shown.

Vendor profiles

ANTHROPIC

+ Affirmative

Claude Opus 5

claude-opus-5

Yes.



Speed	83%
First-token responsiveness	83%
Token economy	32%

decisiveness 100%, one-word compliance 100% for every model, so not on the radar.

Input tokens	25
Output tokens	35
Latency	1.4 s
First token	1.4 s
Throughput	25.3 tok/s
Cost	\$0.002950

Picks an answer and returns it promptly. Conviction and speed.

ANTHROPIC

+ Affirmative

Claude Sonnet 5

claude-sonnet-5

Yes



Speed	87%
First-token responsiveness	94%
Token economy	92%

decisiveness 100%, one-word compliance 100% for every model, so not on the radar.

Input tokens	25
Output tokens	5
Latency	1.2 s
First token	664 ms
Throughput	4.3 tok/s
Cost	\$0.000300

Picks an answer and returns it promptly. Conviction and speed.

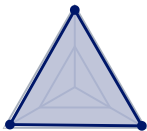
ANTHROPIC

+ Affirmative

Claude
Haiku 4.5

claude-haiku-4-5-
20251001

Yes.



Speed	95%
First-token responsiveness	95%
Token economy	92%

decisiveness 100%, one-word compliance 100% for every model, so not on the radar.

Input tokens	20
Output tokens	5
Latency	645 ms
First token	606 ms
Throughput	7.8 tok/s
Cost	\$0.000135

Picks an answer and returns it promptly. Conviction and speed.

OPENAI

+ Affirmative

GPT-5.6 Sol

gpt-5.6-sol

Yes.



Speed	88%
First-token responsiveness	91%
Token economy	90%

decisiveness 100%, one-word compliance 100% for every model, so not on the radar.

Input tokens	17
Output tokens	6
Latency	1.1 s
First token	906 ms
Throughput	5.4 tok/s
Cost	\$0.000564

Picks an answer and returns it promptly. Conviction and speed.

OPENAI
GPT-5.5
gpt-5.5

+ Affirmative

Yes



Speed	88%
First-token responsiveness	91%
Token economy	62%

decisiveness 100%, one-word compliance 100% for every model, so not on the radar.

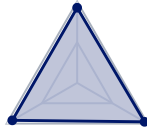
Input tokens	17
Output tokens	20
Latency	1.1 s
First token	855 ms
Throughput	18.8 tok/s
Cost	\$0.002025

Picks an answer and returns it promptly. Conviction and speed.

OPENAI
GPT-5.4 mini
gpt-5.4-mini

+ Affirmative

Yes



Speed	92%
First-token responsiveness	96%
Token economy	92%

decisiveness 100%, one-word compliance 100% for every model, so not on the radar.

Input tokens	17
Output tokens	5
Latency	801 ms
First token	576 ms
Throughput	6.2 tok/s
Cost	\$0.000105

Picks an answer and returns it promptly. Conviction and speed.

XAI

+ Affirmative

Grok 4.6

grok-4.6

Yes



Speed	0%
First-token responsiveness	0%
Token economy	100%

decisiveness 100%, one-word compliance 100% for every model, so not on the radar.

Input tokens	647
Output tokens	1
Latency	6.8 s
First token	6.7 s
Throughput	0.1 tok/s
Cost	\$0.003900

Picks a clear answer but takes its time. Conviction over speed.

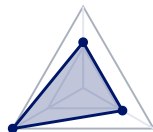
XAI

+ Affirmative

Grok 4.3

grok-4.3

Yes



Speed	56%
First-token responsiveness	56%
Token economy	100%

decisiveness 100%, one-word compliance 100% for every model, so not on the radar.

Input tokens	203
Output tokens	1
Latency	3.2 s
First token	3.1 s
Throughput	0.3 tok/s
Cost	\$0.000768

Picks an answer and returns it promptly. Conviction and speed.

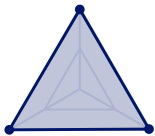
XAI

+ Affirmative

Grok 4.20 (non-reasoning)

grok-4.20-0309-non-reasoning

Yes



Speed	98%
First-token responsiveness	98%
Token economy	100%

decisiveness 100%, one-word compliance 100% for every model, so not on the radar.

Input tokens	195
Output tokens	1
Latency	444 ms
First token	414 ms
Throughput	2.3 tok/s
Cost	\$0.000738

Picks an answer and returns it promptly. Conviction and speed.

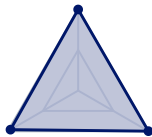
MISTRAL AI

+ Affirmative

Mistral Medium 3.5

mistral-medium-2604

Yes.



Speed	100%
First-token responsiveness	100%
Token economy	96%

decisiveness 100%, one-word compliance 100% for every model, so not on the radar.

Input tokens	27
Output tokens	3
Latency	332 ms
First token	313 ms
Throughput	9 tok/s
Cost	\$0.000189

Picks an answer and returns it promptly. Conviction and speed.

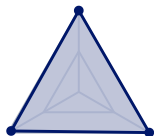
MISTRAL AI

+ Affirmative

Mistral Small 4

mistral-small-2603

Yes



Speed	100%
First-token responsiveness	100%
Token economy	96%

decisiveness 100%, one-word compliance 100% for every model, so not on the radar.

Input tokens	27
Output tokens	3
Latency	311 ms
First token	294 ms
Throughput	8 tok/s
Cost	\$0.000017

Picks an answer and returns it promptly. Conviction and speed.

DEEPSEEK

+ Affirmative

DeepSeek V4 Pro

deepseek-v4-pro

Yes



Speed	78%
First-token responsiveness	78%
Token economy	0%

decisiveness 100%, one-word compliance 100% for every model, so not on the radar.

Input tokens	94
Output tokens	51
Latency	1.7 s
First token	1.7 s
Throughput	29.3 tok/s
Cost	\$0.000483

Picks a clear answer but takes its time. Conviction over speed.

[Source on GitHub](#) · [Methodology](#) · [MIT license](#) · [RSS](#) · [JSON feed](#) · [Built with n-dx](#)

An En Dash Consulting research program of no consequence whatsoever, published weekly. The questions are silly on purpose; the measurement is not. See [About](#).

More from En Dash, free and no API key: [Learn LangGraph](#) · [Learn LangChain](#) · [Learn AgentCore](#)