

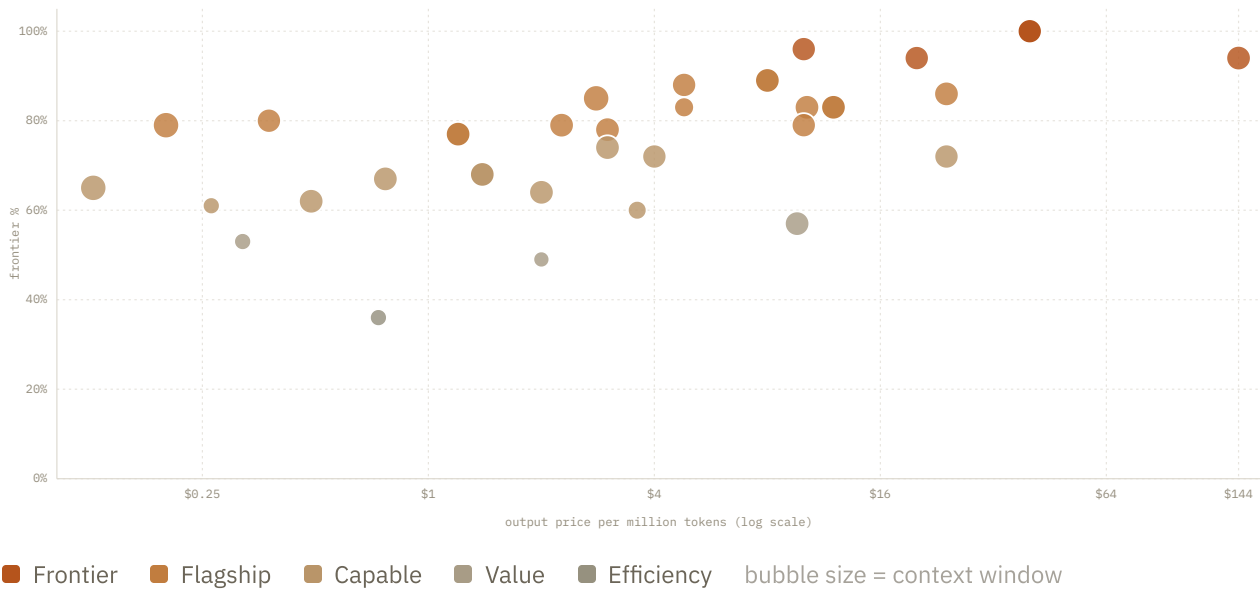
Frontier Index

Every one of the **334 distinct chat-capable model IDs** currently served through Nous Portal's inference API — pulled live from its own `/v1/models` endpoint, not Hermes's curated picker shortlist. **37 of them** carry a published Artificial Analysis Intelligence Index reading; the rest are ranked by real spec data only (price, context, tool-calling, reasoning, modality, architecture) — never a guessed score. Pulled 15 September 2026.

SMARTEST	BEST VALUE	CHEAPEST (PAID)	LONGEST CONTEXT
GPT-6 Ast...	DeepSeek:...	Mistral: M...	SpaceXAI:...
100% of frontier	65% at \$0.13/M out	\$0.024/M out · 7 free-tier models exist too	2M tokens

Price vs. intelligence

the 37 models with a published score · output price, log scale · bubble = context window · hover any point



Search model, provider, or id...

All Frontier Flagship Capable Value Efficiency Unscored 334 of 334 models						
#	MODEL	FRONTIER % ▾	PRICE /M TOK ▾	CONTEXT ▾	CAPABILITIES	
1	● OpenAI: GPT-6 Astra openai/gpt-6-astra	100% AA ● Frontier	\$8.00 in · \$40.00 out	1.1M	✂ Tool calling 🧠 Always-on max/xhigh/ 👁 File+Image	
2	● OpenAI: GPT-6 Astra Pro openai/gpt-6-astra-pro	100% AA ● Frontier	\$8.00 in · \$40.00 out	1.1M	✂ Tool calling 🧠 Always-on max/xhigh/ 👁 File+Image	
3	● Anthropic: Claude Fable 5.1 anthropic/claude-fable-5.1	100% AA ● Frontier	\$8.00 in · \$40.00 out	1M	✂ Tool calling 🧠 Always-on max/xhigh/ 👁 Text+Image	
4	● Claude Opus 5 (batch) anthropic/claude-opus-5:batch live via Nous: batch only	96% AA ● Frontier	\$2.00 in · \$10.00 out	1M	✂ Tool calling 🧠 On by default max/xhigh/ 👁 Text+Image	
5	● Anthropic: Claude Fable 5 (batch) anthropic/claude-fable-5:batch live via Nous: batch only	94% B ● Frontier	\$4.00 in · \$20.00 out	1M	✂ Tool calling 🧠 Always-on max/xhigh/ 👁 Text+Image	

Methodology. "Frontier %" = a model's Artificial Analysis Intelligence Index score ÷ 53 (the confirmed top score in this catalog — Claude Fable 5.1 and GPT-6 Astra are tied there) × 100. Source tags: AA = fetched directly from artificialanalysis.ai (highest confidence, 8 models cross-checked head-to-head); B = BenchLM's leaderboard table, which matched AA-direct within ~1pt on five separate models we verified; B* = BenchLM's number for the GPT-5.6 Sol/Terra/Luna trio was ~25% inflated versus AA-direct on Sol specifically, so Terra/Luna are corrected by that same

ratio rather than used raw; M = ModelGrep, used only where BenchLM had no coverage; E = no independent score published anywhere — not estimated numerically, shown as unscored and positioned by price/lineage only. **297 of 334 models have no published Intelligence Index at all** and are intentionally left unscored rather than guessed.

Scope note. Five current-generation Claude models (Opus 5, Sonnet 5, Haiku 4.5, Fable 5, Opus 4.8) have no synchronous entry in Nous's live catalog right now — only an `async : batch` endpoint, flagged in the table. Thirteen `~org/model-latest` alias pointers were excluded as duplicates of a model already listed under its real id. Tool-calling support is shown per model but is now universal across every non-embedding chat model in this catalog (334/334) — not a differentiator, just verified rather than assumed.

Sources: [Artificial Analysis Intelligence Index v4.3](#) · [AA: Astra vs Sol](#) · [BenchLM](#) · [ModelGrep](#) · [Nous Portal inference-api.nousresearch.com/v1/models](#) (live, unauthenticated, fetched at publish time)