



Anthropic Cuts Fable 5.1 Cache Reads 75% to \$0.25/M — and Breaks Three Prompt Patterns Developers Rely On



Anthropic Cuts Fable 5.1 Cache Reads 75% to \$0.25/M — and Breaks Three Prompt Patterns Developers Rely On

Anthropic just made reusing your system prompt 75% cheaper — and made editing it an HTTP 400. On Fable 5.1, the cheapest prompt is the one you never touch.

TL;DR

On September 1, 2026 Anthropic shipped Claude Fable 5.1 and Mythos 5.1 with base pricing unchanged at \$10/\$50 per million tokens, but prompt cache reads cut from \$1.00 to \$0.25 per million — 0.025× the input rate instead of the usual 0.1×. Three breaking changes come with it: forced `tool_choice` now returns HTTP 400, Fable 5.1 thinking blocks are unreadable by earlier models, and editing any earlier turn — system prompt, tools array, or message history — invalidates the cached reasoning. The economics inverted: a large frozen prefix is now the cheapest architecture, and the mutable, edit-in-place prompting most teams built is both the



Anthropic Cuts Fable 5.1 Cache Reads 75% to \$0.25/M — and Breaks Three Prompt Patterns Developers Rely On

expensive path and the error-producing one.

The news

Anthropic released [Claude Fable 5.1](#) and Claude Mythos 5.1 on September 1, 2026. Base token pricing did not move: \$10 per million input tokens and \$50 per million output tokens for both models, identical to Fable 5. Batch API still runs at half rates — \$5.00 in, \$25.00 out. Cache writes are unchanged at \$12.50 per million for the 5-minute tier and \$20.00 for the 1-hour tier.

What moved is the read side. Prompt cache reads dropped from \$1.00 to \$0.25 per million tokens. Per the [Claude Platform API release notes](#), that puts Fable 5.1 cache reads at 0.025× the base input rate, down from the 0.1× ratio that applies across other Claude models. Anthropic estimates the change cuts typical workload cost by roughly 25%, and highly agentic workloads by up to roughly 45%.

Both models ship with a 1M-token context window by default, 128k max output tokens, and always-on adaptive thinking. Mythos 5.1 is listed as a gated preview variant for vetted defenders and security researchers at the same \$10/\$50 pricing.

\$0.25/M

Fable 5.1 prompt cache reads, down from \$1.00

Anthropic Claude Fable product page, 2026-09-01

0.025×

cache read as a multiple of base input rate (was 0.1× on other Claude models)

Claude Platform API release notes, 2026-09-03

~45%

Anthropic-estimated cost cut on highly agentic workloads

Anthropic Claude Fable product page, 2026-09-01

Alongside the launch, Anthropic introduced Enterprise Frontier Safeguards, which stores monitoring data inside the customer's own AWS, Azure, or Google Cloud environment with customer-managed encryption keys and access policies. That is a procurement story more than an engineering one, but it removes one of the standard objections security teams raise about frontier-model telemetry.



The three breaking changes

The pricing headline traveled further than the compatibility notes, which is backwards. Per Anthropic's ["What's new in Claude Fable 5.1"](#) documentation:

1. Forced tool use is gone. `tool_choice` set to `{"type": "any"}` or `{"type": "tool", "name": ...}` now returns an HTTP 400 error on Fable 5.1. If your agent loop relies on compelling a tool call at a specific step — a router that must emit a structured decision, a validator that must call a schema-checking function — that code path fails immediately on the new model.

2. Thinking blocks move forward, not backward. Fable 5.1 can read thinking blocks produced by earlier Claude models. Earlier models cannot read Fable 5.1's blocks. Any architecture that mixes models in one conversation — cheap model for triage, expensive model for the hard step, back to cheap for summarization — now has a one-way valve in it.

3. Editing history invalidates reasoning. Modifying the system prompt, the tools array, or any earlier message invalidates the Fable 5.1 thinking block on the next request. And per [reporting on the launch](#), for API accounts created on or after August 31, 2026, editing history before a thinking block returns a 400 error rather than silently dropping the block.

!

Watch out That August 31, 2026 account cutoff means two teams running identical code get different behavior. An older account degrades quietly — the thinking block disappears, quality drops, nobody gets paged. A newer account throws a 400 and you find out in seconds. If your staging environment predates the cutoff and production doesn't, your integration tests are lying to you.

Why it matters: the cheapest prompt is the one you never touch

The three breaking changes look unrelated. They aren't. Read together with the pricing move, they describe a single architectural preference: **Anthropic wants your context to be append-only.**



Anthropic Cuts Fable 5.1 Cache Reads 75% to \$0.25/M — and Breaks Three Prompt Patterns Developers Rely On

Cache reads at 0.025× base input make a large frozen prefix nearly free to reuse. A 200k-token system prompt plus tools plus retrieved corpus costs \$2.00 per call at base input rates. Cached, it costs five cents. The cache write costs \$12.50 per million, so a 200k prefix costs \$2.50 to write once — and you amortize that across every subsequent call in the TTL window. The break-even is roughly the second call.

Then the invalidation rule closes the loop. Touch that prefix and you don't just pay a cache write again — on new accounts, you get an error. The pricing rewards immutability and the API enforces it.

Anthropic didn't cut a price. They repriced a design pattern and deprecated its alternative in the same release.

Who wins: teams that already treat the system prompt as a build artifact — versioned, compiled, deployed, immutable at runtime. Teams running long agentic sessions where the same tool definitions and instructions ride along for hundreds of turns. That is exactly the population Anthropic's ~45% agentic-workload estimate is aimed at.

Who loses: teams whose prompt logic mutates per request. Dynamic tool arrays assembled from user permissions. System prompts with the current timestamp interpolated in. Conversation compaction that rewrites earlier turns to save tokens. Every one of those patterns now busts cache on each call, and on post-cutoff accounts several of them throw 400s.

i

The timestamp trap Injecting a live clock value into the system prompt is one of the most common prompt patterns in production agents. Under the new economics it is also one of the most expensive: it guarantees a cache miss on every single call, turning a five-cent read into a two-dollar input charge on a 200k prefix. Move volatile values to the end of the message array, not the top.

The benchmarks, and what they actually support

Fable 5.1's capability jump is real and it is concentrated in agentic and terminal



Anthropic Cuts Fable 5.1 Cache Reads 75% to \$0.25/M — and Breaks Three Prompt Patterns Developers Rely On

work.

Benchmark	Fable 5.1	Fable 5	Opus 5	GPT-5.6	Sol
Terminal-Bench-Science 0.1	52.6%	24.7%	29.0%	22.4%	
Terminal-Bench 4.0	55.8%	42.0%	52.3%	37.3%	
AutomationBench	31.4%	17.1%	26.9%	19.6%	
CursorBench 3.2.0	73.4%	70.5%	70.0%	67.2%	

The Terminal-Bench-Science number is the eye-catcher — 52.6% against Fable 5’s 24.7%, a 2.1× jump, reported by [R&D World](#) and tabulated by [Apidog](#). AutomationBench moved 17.1% → 31.4%, an 83% relative gain, per [Vellum](#).

Note the shape of the spread. On CursorBench — code editing in an IDE harness — the gap over Fable 5 is 2.9 points. On Terminal-Bench-Science it is 27.9 points. The gains cluster where the model runs long autonomous loops with tools, not where it does single-shot edits. That is consistent with the pricing move: both are optimized for the same workload.

Treat all of these numbers with the caveat that harness design moves scores as much as model weights do. I’ve written before about [why the same scaffold scores 65% or 74% on SWE-bench](#) — the scaffold, retry policy, and tool surface account for a swing wide enough to swallow most vendor-reported deltas. A 27.9-point gap on Terminal-Bench-Science is too large to be harness noise. A 2.9-point gap on CursorBench is not.

My take: this is not a discount

My take

Most coverage framed this as a price cut. It is a migration cost transferred to you, partially rebated in tokens. The 75% cache read reduction is worth real money only if your context is already immutable — and if it isn’t, the rewrite to make it immutable will cost more engineering hours than the first year of savings. Anthropic is buying a behavior change with a discount, and the discount is priced to make the change look voluntary.

The genuinely under-discussed part: the invalidation rule and the 1M-token context



Anthropic Cuts Fable 5.1 Cache Reads 75% to \$0.25/M — and Breaks Three Prompt Patterns Developers Rely On

window pull in opposite directions. A million tokens invites you to stuff everything into context. But the standard technique for surviving long sessions — periodically compacting or rewriting earlier turns to control cost and drift — is now the thing that invalidates your cached reasoning, and on post-cutoff accounts, errors out. You get a bigger window and less freedom to manage what's inside it.

That matters because context length is not free in quality terms either. My own analysis across [172 billion tokens found every model tested fabricating above 10% at 200k context](#). The cheap-cache incentive pushes toward larger frozen prefixes; the fabrication curve pushes toward smaller working sets. Those are not the same optimum, and the pricing sheet only argues one side.

The claim

Fable 5.1 makes long-context agents 25–45% cheaper, so put more into context.

The reality

Anthropic's own ~25%/~45% figures are estimates for typical and highly agentic workloads respectively, and they assume cache hits. Every mutation of the system prompt, tools array, or earlier messages forfeits the read price and, on accounts created on or after August 31, 2026, returns a 400. Cheap context is conditional on frozen context.

What to do this week

If you run Fable on production traffic, the migration is small but it is not zero — and two of the three breaking changes fail loudly rather than silently, which is good news for anyone who tests before shipping.

Grep for tool_choice

Any occurrence of `{"type": "any"}` or `{"type": "tool", "name": ...}` is a guaranteed HTTP 400 on Fable 5.1. Replace forced tool calls with instruction-level constraints plus output validation. If you were forcing tool use to guarantee structured output, you now need a schema check and a retry, not an API flag.



Audit what changes between turns

Diff your outgoing request payloads across consecutive calls in a real session. Anything that moves in the system prompt, the tools array, or prior messages is a cache miss and an invalidation. Timestamps, per-user tool subsets, and injected retrieval results are the usual offenders.

Push volatility to the tail

Restructure so the frozen prefix — instructions, tool definitions, stable reference material — sits at the top and never changes, and everything variable appends at the end. This is the single change that converts the 0.025× read rate from a headline into a line-item.

Check your account creation date

Accounts created on or after August 31, 2026 get a 400 on stale thinking blocks. Older accounts drop the block silently. Confirm which behavior your staging and production accounts have — if they differ, your test suite is validating the wrong failure mode.

Map your multi-model handoffs

Thinking blocks are forward-compatible only. Any pipeline that routes from Fable 5.1 back down to an earlier Claude model needs the reasoning stripped at the boundary. Decide deliberately whether you strip it or restructure the pipeline so the expensive model runs last.

✓

The fix Treat the system prompt and tools array as a compiled, versioned build artifact with a content hash. Deploy changes; never patch them at request time. This single discipline satisfies the cache economics, the invalidation rule, and the 400-error behavior simultaneously — and it makes prompt regressions traceable to a commit.



Where this goes

I expect the 0.025× ratio to spread. Anthropic priced cache reads on Fable 5.1 at a quarter of the ratio used on its other models. That is not a stable equilibrium across a product line. Either it becomes the house rate or Fable 5.1 becomes the model everyone routes agentic traffic to for reasons that have nothing to do with its Terminal-Bench score.

I expect prompt immutability to become a framework feature, not a discipline. Right now, keeping a prefix frozen is something your team has to remember to do. Within six to twelve months I expect the major agent frameworks to ship append-only context objects that make mutation of the prefix a compile-time error rather than a runtime billing surprise. The pricing incentive is strong enough and the failure mode is loud enough that someone will build it.

What I am not going to claim: that the cache TTL changed. Anthropic's Fable 5.1 documentation does not state a changed time-to-live. The 5-minute and 1-hour tiers and their write prices are what they were. The pricing changed; the TTL policy did not. If your economics depend on holding a cache warm across a longer idle window, nothing in this release helps you.

What works

- Cache reads at \$0.25/M make large frozen prefixes economically dominant — roughly break-even on the second call.
- Real capability jump where it counts for agents: 52.6% vs 24.7% on Terminal-Bench-Science, 31.4% vs 17.1% on AutomationBench.
- Two of three breaking changes fail loudly (HTTP 400), so migration bugs surface in testing rather than in production quality drift.
- Enterprise Frontier Safeguards keeps monitoring data in the customer's own cloud with customer-managed keys.

Rough edges

- Forced `tool_choice` removal breaks router and structured-output patterns with no drop-in replacement.
- Behavior splits on account creation date — pre- and post-August 31, 2026 accounts fail differently on the same code.
- Thinking blocks are forward-compatible only, which constrains cheap-



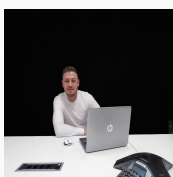
Anthropic Cuts Fable 5.1 Cache Reads 75% to \$0.25/M — and Breaks Three Prompt Patterns Developers Rely On

model/expensive-model pipelines.

- The 1M window invites large context while the invalidation rule penalizes compaction — and long context carries its own fabrication cost.

Bottom line

The 75% cache read cut is only a discount if your context is already immutable; otherwise it is a bill for the refactor that makes it so. Freeze the prefix, append the volatility, kill your forced `tool_choice` calls, and check which side of August 31, 2026 your account falls on. If you want a second pair of eyes on your agent's context architecture before the migration, [Book a call →](#)



WRITTEN BY

Artur Markus

Software engineer in Basel building AI infrastructure and agentic systems, with a background in pharmaceutical automation, GxP compliance, and validation. I help teams turn AI from a chatbot into a reliable operational layer.

Book a meeting →