

GLM-5.3 vs GLM-5.3-Flash: Post-Training and Deployment Guide

IntuitionLabs.ai · Adrien Laurent · 2026-09-05 · glm-5.3 · glm-5.3-flash · zhipu ai · z.ai · open weight llm · mixture of experts · llm pricing comparison · ai model licensing · multimodal llm · llm deployment guide



Executive Summary

GLM-5.3 and **GLM-5.3-Flash**, released by Zhipu AI under its international **Z.ai** brand on August 14 and August 26, 2026 respectively ([the-decoder.com](#)), share a family name but not a foundation. GLM-5.3 reuses GLM-5.2's base model and improves purely through an extended post-training phase, arriving at **744 billion total parameters with 40 billion active** in a mixture-of-experts design and a **bespoke commercial license** that gates large Model-as-a-Service operators behind a Z.AI security review once combined revenue with affiliates exceeds \$10 billion over any 12 months. GLM-5.3-Flash, by contrast, starts from an entirely new base model, adds native multimodal input (video, image, text, and file, though output remains text-only), and ships at **320 billion total and 18 billion active parameters** under a plain, unmodified **MIT license** ([benchgen.com](#)).

Pricing is the sharpest divergence. Z.ai's own API rates list GLM-5.3-Flash at **\$0.15 per million input tokens and \$0.50 per million output tokens**, against GLM-5.3's **\$1.40 and \$4.40**, a roughly nine-fold gap that is broadly consistent with Z.ai's own characterization of Flash as available at one-tenth the price of the prior generation. A

limited-time 50% launch discount on Flash pricing, corroborated independently by third-party aggregator OpenRouter, expires **September 9, 2026** (openrouter.ai). Both models are also bundled into the **GLM Coding Plan** subscription, which starts at \$18 per month, though GLM-5.3 consumes roughly three times more of a subscription's credit quota per token than GLM-5.3-Flash does.

On benchmarks, Z.ai's own comparison tables show both models improving substantially over GLM-5.2 on agentic coding tasks, with GLM-5.3 posting the larger gains; the one figure in this comparison independently sourced from a third party, Artificial Analysis's Intelligence Index, scored GLM-5.3-Flash at **57, ranking it 10th globally** at the time of testing (scmp.com).

For buyers, the choice reduces to license exposure, modality needs, and budget: GLM-5.3 suits text-only agentic workloads at large organizations comfortably clear of its revenue-gated license threshold, while GLM-5.3-Flash's unrestricted MIT terms, multimodal input support, and lower cost make it the structurally simpler default for most production deployments, including at regulated enterprises that must document their tooling choices before adoption.

Introduction and Background

Zhipu AI, which operates internationally under the brand **Z.ai**, released two large language models in quick succession in August 2026: the flagship **GLM-5.3** on August 14 (the-decoder.com), and **GLM-5.3-Flash** twelve days later, on August 26 (z.ai). Despite the near-identical name, the two models are built on different foundations, target different deployment budgets, and carry different licenses, which makes "GLM-5.3 vs GLM-5.3-Flash" a genuinely substantive comparison rather than a simple size variant question.

Z.ai is the international brand of Zhipu AI, a company founded in Beijing in 2019 as a Tsinghua University spin-out (en.wikipedia.org). The company rebranded itself as Z.ai internationally in July 2025 (en.wikipedia.org) and completed an initial public offering on the Hong Kong Stock Exchange on January 8, 2026 (en.wikipedia.org). Its GLM (General Language Model) series progressed through GLM-4.5 and GLM-4.6 before GLM-5.2 **was released with a one-million token context window** on June 16, 2026 (en.wikipedia.org), setting up the base model that GLM-5.3 would later build on with post-training alone.

This report treats GLM-5.3 and GLM-5.3-Flash as parallel, independently deployable models rather than a single tier ladder, since their architectures, licenses, and target use cases genuinely diverge. It documents what changed technically between GLM-5.2, GLM-5.3, and GLM-5.3-Flash; what each model costs to call through Z.ai's own API and through the GLM Coding Plan subscription; what benchmark evidence exists and how it is sourced; and what the licensing and deployment differences mean for teams evaluating either model, including regulated enterprises that must document their AI tooling choices. As with any comparison of open-weight and API-hosted models, buyers evaluating this class of model, including regulated-industry teams that must document AI tooling decisions for internal governance, benefit from treating vendor benchmark claims and self-hosting requirements as separate questions from list pricing; IntuitionLabs, a life-sciences and AI consultancy, frames this as building a **"governed information"** layer before adoption decisions are made, connecting a model's output to **"authoritative enterprise sources with identity, permissions, retrieval, citations, evaluation, and accountable operation"** rather than adopting a model on benchmark scores alone (intuitionlabs.ai). All prices, benchmark figures, and availability statements below are anchored to their observation date of **September 5, 2026**, because Z.ai's own pricing page carries a promotional window that expires days after that date.

GLM-5.3: Capabilities, Adoption, and Limitations

Architecture and Post-Training

GLM-5.3's defining technical characteristic is what Z.ai chose not to change. The model **uses the same base model as GLM-5.2**, meaning every capability gain in this release comes from an extended post-training phase rather than a new pretraining run (huggingface.co). Z.ai's own account of the release opens with that framing directly: post-training scale-up, not architecture change, produced GLM-5.3 (interconnects.ai). The post-training stack carries over techniques introduced with GLM-5.2, including **IndexShare** for efficient long-context processing and **SAO** (a reinforcement-learning method for long-horizon agentic tasks) (z.ai).

Structurally, GLM-5.3 is a **mixture-of-experts (MoE)** model with **744 billion total parameters and 40 billion active parameters**, distributed as native FP8 weights (raw.githubusercontent.com). That scale reflects the broader GLM-5 architecture, which expanded from GLM-4.5's 355 billion total / 32 billion active parameter design to 744 billion / 40 billion active for the GLM-5 generation (raw.githubusercontent.com). GLM-5.3 supports a **1-million-token context window** with text-only inputs (docs.z.ai), a figure independently listed by third-party host Cloudflare as pairing that context length with reasoning, function calling, and structured outputs (developers.cloudflare.com).

Licensing is a departure from Z.ai's earlier practice: the Hugging Face model card lists the license as **"other"** with a custom license name, **glm-5.3**, rather than a standard permissive license (huggingface.co). That bespoke license requires large commercial operators, specifically Model-as-a-Service providers whose aggregate revenue with affiliates exceeds **\$10 billion** over any trailing 12 months, to pass a Z.AI security review before commercial use (huggingface.co) (huggingface.co). Independent commentary characterizes it as **"a bespoke, vendor-named licence written for this one model"** rather than a reused open-source template (digitalapplied.com), a distinction with direct enterprise-deployment consequences discussed in the licensing section below. The separate GitHub repository containing inference and serving code, as opposed to the model weights themselves, is released under the standard **Apache-2.0** license (github.com).

Adoption and Availability

GLM-5.3 is available through the **GLM Coding Plan**, and Z.ai's documentation lists Model API endpoints for the OpenAI Chat Completion, OpenAI Response, and Anthropic Message protocols (docs.z.ai). Z.ai committed to releasing the open weights approximately **two weeks after launch**, once safety evaluation and hardening were complete, a timeline corroborated independently (z.ai) (the-decoder.com). Self-hosted deployment is documented against multiple open-source inference engines, including **SGLang, vLLM, TokenSpeed, Transformers, KTransformers, and Unsloth**, with a published Docker-based serving path for SGLang (huggingface.co).

Third-party hosting exists alongside Z.ai's documented Model API: Cloudflare Workers AI lists GLM-5.3 at **\$1.40 per million input tokens and \$4.40 per million output tokens**, with cached input at \$0.26 per million (developers.cloudflare.com), figures that match Z.ai's own published API pricing (see the Data Analysis section). OpenRouter, a separate third-party aggregator, lists a different rate of **\$1.15 in / \$3.50 out per million tokens** (openrouter.ai), a discrepancy worth flagging: third-party hosts negotiate independent terms and are not obligated to match a vendor's list price, so buyers comparing quotes should confirm current pricing directly with whichever host they intend to use rather than assuming uniformity across providers.

Strengths and Limitations

On coding and agentic benchmarks, Z.ai's own comparison table reports substantial gains over GLM-5.2 (figures given in the Performance and Benchmarks section below), alongside a claimed **50% improvement** on Z.ai's in-house Code Bench (huggingface.co). These are vendor-reported figures from Z.ai's own comparison methodology; one benchmark in the same table, **GDPval-AA v2**, is explicitly attributed to an independent evaluator rather than Z.ai itself, with the model card stating models on that benchmark **"are evaluated by Artificial Analysis"** (huggingface.co), a useful marker of which numbers in Z.ai's own materials carry third-party verification and which do not.

Z.ai also documents an emergent capability that is unusual for a vendor to disclose in benchmark terms: coding-focused post-training scale-up produced cyber-relevant reasoning that, according to the vendor, reasons across multiple stages of a vulnerability-exploitation chain and forms coherent multi-step plans. In a disclosed vulnerability-research collaboration, this capability **identified 2,436 vulnerabilities across 269 open-source projects, including 1,097 of medium-to-high severity** (z.ai), a defensive security-research use documented by the vendor rather than an incident involving any named deployment. A documented limitation is a behavioral regression from earlier GLM releases: **"disabling thinking is no longer supported by GLM-5.3,"** meaning the model's extended-reasoning mode cannot be turned off for latency-sensitive use cases the way it could in prior versions (z.ai). Community reception on Hacker News, a discussion venue rather than a benchmark source, was broadly favorable, with one widely upvoted comment describing GLM-5.3 as **"probably the sweet spot open weights model if you want to go beyond deepseek flash"** and another noting it **"has been able to tackle all the random hard problems I've thrown at it"** (news.ycombinator.com) (news.ycombinator.com); this is developer sentiment, not a controlled evaluation, and is labeled as such.

GLM-5.3-Flash: Capabilities, Adoption, and Limitations

Architecture and Multimodal Capabilities

GLM-5.3-Flash, released **August 26, 2026**, is described by Z.ai as **"the first natively multimodal model in the GLM-5 series"** (z.ai). Unlike GLM-5.3, which reuses GLM-5.2's base model, GLM-5.3-Flash **starts from a newly trained base model**, giving Z.ai room to change the underlying architecture rather than only the post-training recipe, a distinction carried in the GLM-5 series' own reference documentation (raw.githubusercontent.com). It is a smaller mixture-of-experts model than its flagship sibling, at **320 billion total parameters and 18 billion active parameters** (raw.githubusercontent.com), routed across 288 experts with 8 activated per token according to its published configuration file (huggingface.co).

"Multimodal" here refers specifically to input understanding rather than generation: documented input modalities are **video, image, text, and file**, while output remains **text-only**, with a 1-million-token context length and a 128,000-token maximum output (docs.z.ai). Cloudflare's independent listing corroborates the exact context figure, **1,048,576 tokens**, alongside vision, function-calling, and reasoning support (developers.cloudflare.com). On the efficiency side, Z.ai reports that GLM-5.3-Flash uses **3.0 times less attention compute and a 4.4 times smaller key-value cache** than GLM-5.3 at comparable context lengths (z.ai), a serving-cost difference that matters more for deployment planning than raw parameter counts do.

A second architectural distinction concerns licensing rather than model weights: where GLM-5.3 ships under a bespoke license, GLM-5.3-Flash's published license file on Hugging Face is the standard, unmodified **MIT License** text, beginning **"Permission is hereby granted, free of charge, to any person obtaining a copy of this software"** (huggingface.co), a characterization repeated by independent model trackers as **"open weights (MIT), free to download and self-host"** (benchgen.com). That difference, no revenue-gated security review clause and no custom terms, is the single clearest deployment-planning distinction between the two models for any organization evaluating self-hosting.

Adoption and Availability

Deployment options are broader than GLM-5.3's at launch. GLM-5.3-Flash is directly hosted on **Cloudflare Workers AI** under the model identifier `@cf/zai-org/glm-5.3-flash` (developers.cloudflare.com), is self-hostable through the same set of engines documented for GLM-5.3, **SGLang, vLLM, TokenSpeed, Transformers, KTransformers, and Unsloth** (huggingface.co), and is fully available on the GLM Coding Plan. It is also bundled into the same **GLM Coding Plan** subscription used for GLM-5.3, detailed in the Data Analysis section below.

Strengths and Limitations

Z.ai's own comparison table reports GLM-5.3-Flash improving markedly over GLM-5.2 on agentic coding tasks (figures given in the Performance and Benchmarks section below), gains that place a fraction-of-the-size model close to, and on some measures above, its full-scale predecessor. Independent evaluator Artificial Analysis reportedly scored the model **57 on its Intelligence Index, ranking it 10th globally** at the time of testing (scmp.com), a result that carries more evidentiary weight than Z.ai's in-house tables because it comes from a third-party benchmarking organization rather than the model's own vendor.

The clearest limitation is scope rather than raw capability: multimodal support is input-only. Teams needing image or video generation, or audio input or output, will not find that in GLM-5.3-Flash's documented feature set, only text, image, video, and file understanding feeding into text responses. The model's smaller active-parameter footprint (18 billion versus GLM-5.3's 40 billion) also means it is not presented by Z.ai as a strict capability upgrade over the flagship model; rather, it is positioned as a cost-and-latency-optimized alternative that approaches, without fully matching, GLM-5.3-class performance on the hardest agentic tasks. Independent commentary corroborates this framing directly, describing the flagship GLM-5.3 as **"the proprietary GLM-5.3 flagship (same generation, API-only)"** that **"scores higher"** on the hardest benchmarks even as GLM-5.3-Flash undercuts it sharply on price (benchgen.com).

Feature Comparison

Table 1 below summarizes the documented architectural, licensing, and modality differences between the two models as of the observation date.

Dimension	GLM-5.3	GLM-5.3-Flash
Release date	August 14, 2026	August 26, 2026

Dimension	GLM-5.3	GLM-5.3-Flash
Base model	Reuses GLM-5.2's base model; gains from post-training only	Newly trained base model with revised architecture
Total / active parameters (MoE)	744B total / 40B active	320B total / 18B active
Context window	1,048,576 tokens	1,048,576 tokens
Input modalities	Text only	Video, image, text, file
Output modality	Text	Text only (no image/video/audio generation)
Reasoning toggle	Cannot be disabled	Cannot be disabled; <code>thinking.type</code> supports only enabled
License	Bespoke "glm-5.3" license with revenue-gated commercial security review	Standard, unmodified MIT License
First-party API status (as of Sep 5, 2026)	Documented Model API endpoints plus Coding Plan	Documented Model API plus Coding Plan and Cloudflare Workers AI hosting
Positioning	Flagship, highest-capability agentic/coding model	Cost- and latency-optimized, multimodal-input alternative

Table 1 consolidates figures cited individually in the sections above; see those sections for the underlying source for each row. The table makes the central tradeoff visible in one place: GLM-5.3 offers the larger active-parameter budget, the more mature agentic benchmark record, and a locked-in reasoning mode, at the cost of a more restrictive commercial license and documented Model API endpoints. GLM-5.3-Flash trades roughly half the total parameters and a smaller active-parameter footprint for native multimodal input understanding, a permissive MIT license with no revenue-based gating.

Performance and Benchmarks

Method note: the benchmark figures in this report were compiled from the vendors' and independent evaluators' published pages fetched during this research pass, not from a live head-to-head run against either model's API; no such API access was available in this research environment. Every figure below is labeled by its source type, either "vendor-reported" (from Z.ai's own comparison tables) or "independent" (from a named third-party evaluator such as Artificial Analysis), so that a reader can judge evidentiary weight; none should be read as an estimate or as this report's own measurement.

On Z.ai's own benchmark table, GLM-5.3 shows the larger gains against its GLM-5.2 predecessor on the hardest agentic tasks: a jump from **46.2 to 66.9 on DeepSWE v1.1** and from **23.8 to 28.5 on Agents' Last Exam** ([z.ai](#)), both vendor-reported. GLM-5.3-Flash's improvement over the same GLM-5.2 baseline is smaller in absolute terms on DeepSWE v1.1 (**46.2 to 63.4**) but larger proportionally on AutomationBench (**26.2 to 48.8**) ([z.ai](#)), also vendor-reported, using a comparable methodology since both figures come from the same vendor's comparison exercise against the same GLM-5.2 baseline, which makes the two model-to-predecessor deltas reasonably comparable to each other even though neither is independently audited.

The one figure in this comparison with clear third-party provenance is Artificial Analysis's Intelligence Index score for GLM-5.3-Flash: a reported **57**, placing the model **10th globally** among the models that evaluator tracks (scmp.com). No equivalent independently sourced Intelligence Index figure for GLM-5.3 itself was located during this research pass, and no independent, apples-to-apples benchmark run pitting GLM-5.3 directly against GLM-5.3-Flash on identical prompts was found in any fetched source; readers should treat the vendor's side-by-side comparison table, which does place both models against the shared GLM-5.2 baseline, as the closest available like-for-like comparison, and treat any claim that one model strictly dominates the other on all tasks with appropriate caution.

Data Analysis and Evidence

Pricing is where the two models diverge most sharply in absolute terms. *Table 2* below reproduces Z.ai's own published per-token API rates, including the limited-time promotional discount active at the time of research.

Model	Input (list)	Input (promo, through Sep 9 2026)	Cached input (list)	Output (list)	Output (promo)
GLM-5.3	\$1.40 / M tokens (developers.cloudflare.com)	Not discounted	\$0.26 / M	\$4.40 / M	Not discounted
GLM-5.3-Flash	\$0.15 / M tokens (docs.z.ai)	\$0.075 / M	\$0.03 / M	\$0.50 / M (docs.z.ai)	\$0.25 / M

The promotional pricing for GLM-5.3-Flash **"ends at 24:00 on September 9, 2026 (UTC+8, Singapore time),"** per Z.ai's own pricing documentation (docs.z.ai), a window independently corroborated by OpenRouter's listing of the same discount ending **"September 9, 2026 at 16:00 UTC"** (openrouter.ai). Even at list price, GLM-5.3-Flash costs roughly one-ninth of GLM-5.3 on input tokens (\$0.15 versus \$1.40) and roughly one-ninth on output (\$0.50 versus \$4.40), consistent with Z.ai's own characterization of the smaller model as available **"at one-tenth the price"** of its predecessor generation (huggingface.co). Separately, Z.ai's own documentation frames GLM-5.3-Flash as **"matching Claude Opus 4.8 in intelligence score at just 1/40 the price"** (docs.z.ai), a comparison against a third-party competitor model that should be read as a vendor claim pending independent replication, not as an audited result.

For teams that prefer a subscription model over metered API billing, the **GLM Coding Plan** bundles both models. It starts at **\$18 per month for the Lite tier**, with Pro and Max tiers above it, and all tiers support both GLM-5.3 and GLM-5.3-Flash (docs.z.ai). *Table 3* shows the documented credit quotas across tiers; official monthly prices for the Pro and Max tiers were not confirmed from a directly fetched Z.ai page during this research pass and are left blank rather than estimated.

Plan	Monthly price	5-hour credit quota	Weekly credit quota
Lite	\$18	2,000 (docs.z.ai)	10,000
Pro	Not disclosed in cited docs	12,000	60,000
Max	Not disclosed in cited docs	28,000	140,000

Within the Coding Plan's credit system, the two models are not weighted equally. Documented multipliers show GLM-5.3 consuming **6.9 credits per input token unit, 1.7 for cached input, and 24 for output**, against GLM-5.3-Flash's **2.3, 0.56, and 8** respectively, meaning GLM-5.3 draws down a subscription's quota roughly three times faster than GLM-5.3-Flash for equivalent usage, and usage outside peak hours is billed at half the standard credit rate on either model (docs.z.ai). For teams sizing a subscription rather than metered usage, this multiplier gap is arguably more decision-relevant than the underlying per-token API prices, since it directly determines how many requests a fixed monthly quota will absorb.

Zooming out, the broader competitive context helps explain why Z.ai is pricing this aggressively. **Chinese open-weight frontier models** more generally are positioned as substantially cheaper than comparable Western closed models: one industry analysis found that DeepSeek's V3.2 model costs **roughly 19 times less on input and 66 times less on output** than Claude Opus 4.6 (inferencehub.org), while noting that on one leading leaderboard, **the top 13 spots are all Western models (Anthropic, Google, xAI, OpenAI) (inferencehub.org)**, even as **Alibaba's Qwen family alone has surpassed Meta's Llama in cumulative Hugging Face downloads (inferencehub.org)**. Z.ai's GLM-5.3 and GLM-5.3-Flash pricing sits squarely inside that pattern of open-weight Chinese models undercutting closed Western competitors on cost while trailing on top-end leaderboard rank.

Implications and Future Directions

The practical decision most teams face is not "which model is better" in the abstract but "which model fits a specific workload's cost, license, and modality constraints." For pure text-based agentic coding work where neither the licensee nor an affiliate operates a defined Model-as-a-Service business and their combined revenue exceeds \$10 billion over any consecutive 12 months, GLM-5.3's larger active-parameter budget and stronger agentic benchmark deltas make it the more capable option, at roughly nine times the per-token cost of GLM-5.3-Flash. The license excludes end-user products with model capabilities solely embedded within specific features or harnesses, and mere relaying of requests to models hosted by others. For workloads that need to read screenshots, diagrams, or short video alongside text, or that prefer GLM-5.3-Flash's native multimodal input, or that need a genuinely unrestricted MIT license for redistribution, GLM-5.3-Flash is the structurally better fit regardless of the modest benchmark gap.

The bespoke "glm-5.3" license is worth particular attention for any organization evaluating self-hosting at scale. The review condition applies only where the licensee or an affiliate operates the license's defined Model-as-a-Service business and their combined revenue exceeds \$10 billion over any consecutive 12 months, a detail easy to miss when comparing headline "open weights" claims across models. Organizations building internal evaluation processes for this kind of decision, including regulated industries such as life sciences that must document tooling choices for internal governance, are generally better served by treating license terms, self-hosting requirements, and benchmark provenance as three separate checklist items rather than folding them into a single "is it open source" question. IntuitionLabs, a life-sciences and AI consultancy that is itself a Veeva X-Pages partner rather than a model vendor, frames this kind of evaluation as building organizational capability **"one department at a time," with "governed information, specialist implementation, role-based adoption, and measured results"** preceding any broad rollout decision (intuitionlabs.ai), a framing that applies equally well to choosing between an open-weight coding model and a restrictively licensed flagship. Readers interested in a parallel case study of an open-weight mixture-of-experts model with a different licensing posture, Mistral Large 3's Apache-2.0 release, can consult IntuitionLabs' earlier explainer on that model for contrast (intuitionlabs.ai).

Looking forward, buyers should monitor Z.ai's documentation for API and license changes rather than treating API availability as a differentiator between the two models.

Frequently Asked Questions (FAQs)

What is the core technical difference between GLM-5.3 and GLM-5.3-Flash? GLM-5.3 reuses GLM-5.2's base model and improves purely through post-training (huggingface.co), while GLM-5.3-Flash starts from a newly trained base model with a different architecture and native multimodal input support (raw.githubusercontent.com).

Is GLM-5.3-Flash's multimodality more than image input? As documented in the Architecture and Multimodal Capabilities section above, it accepts video, image, text, and file input but produces text-only output; it does not generate images, video, or audio.

Which model is cheaper to run, and by how much? At list price, GLM-5.3-Flash costs \$0.15 per million input tokens and \$0.50 per million output tokens, against GLM-5.3's \$1.40 and \$4.40 respectively, roughly a nine-fold difference (see Table 2 above for the sourced figures and the promotional discount window).

Can both models be self-hosted? GLM-5.3-Flash is released under a standard MIT license and can be self-hosted without a commercial-scale review requirement (huggingface.co). GLM-5.3 is released under a custom license whose review condition applies when the licensee or an affiliate operates the license's defined Model-as-a-Service business and their combined revenue exceeds \$10 billion over any consecutive 12 months; its stated exclusions are detailed in the GLM-5.3 Architecture and Post-Training section above. Both support SGLang and vLLM serving (huggingface.co).

Was GLM-5.3's cyber capability claim independently verified? No. The 2,436-vulnerability figure discussed above in GLM-5.3's Strengths and Limitations section comes from Z.ai's own disclosed research collaboration; no independently audited replication of that specific figure was located during this research pass, so it should be treated as vendor-reported rather than independently confirmed.

Which model suits enterprise deployment better? It depends on the licensing threshold and modality needs: GLM-5.3-Flash's MIT license and lower per-token cost suit most deployments, while GLM-5.3 suits text-only agentic workloads where its larger active-parameter budget's benchmark edge outweighs its higher cost and more restrictive license, as summarized in Table 1 above.

Conclusion

GLM-5.3 and GLM-5.3-Flash, released twelve days apart in August 2026 by Zhipu AI's international Z.ai brand, are best understood as siblings built for different jobs rather than as a large model and a compressed copy of it. GLM-5.3 keeps GLM-5.2's base model and pushes post-training further, arriving at 744 billion total and 40 billion active parameters, the stronger agentic benchmark record of the two, and a bespoke commercial license whose review condition applies when the licensee or an affiliate operates the license's defined Model-as-a-Service business and their combined revenue exceeds \$10 billion over any consecutive 12 months. GLM-5.3-Flash starts from a new base model entirely, adds native multimodal input understanding, ships at roughly one-ninth the per-token price under a standard MIT license.

For most production teams, the decision comes down to three questions this report has tried to answer with dated, sourced figures rather than adjectives: does the workload need image or video input, does the deploying organization sit anywhere near GLM-5.3's revenue-gated licensing threshold, and does the cost gap, roughly nine-fold on list pricing, matter more than the benchmark gap on the hardest agentic tasks. None of those questions has a universal answer, but each now has a documented one.

For informational purposes. Verify critical medical, legal, financial, regulatory, and technical information with qualified professionals and primary sources.