

Back to Models

Z.AI: GLM 5.3 Flash

z-ai/glm-5.3-flash

Compare API Request Chat

GLM-5.3-Flash is a native multimodal model from Z.ai. It is suited for efficient coding and long-horizon agent tasks. Its hybrid sparse and linear attention architecture maintains accurate long-context behavior while reducing compute overhead.

By: Z.ai Input Type: Output Type: Publish time: 2026-08-26

Providers

Route requests across multiple providers. Copy a provider slug to set your preference.

Provider	Cache Hit Rate	Input	Output	Cache read	Context	Latency	Throughput
BigModel	94.3%	\$0.075 / M tokens	\$0.25 / M tokens	Read: 0.015 / M tokens Write: / M tokens	1M	11.8s	29.4tps

Uptime

Direct request success rate on AI Gateway and per-provider.



ZenMux Announces GLM 5.3 Flash Model Routing Availability

September 02, 2026

Singapore - September 02, 2026 -

ZenMux announced that Z.AI GLM 5.3 Flash is now available through its AI model routing platform, giving developers and enterprise teams a new option for multimodal, coding, and agent-oriented workloads. The model page for ZenMux GLM 5.3 Flash is now live in the company's catalog, presenting provider, pricing, context, latency, throughput, and uptime information in a single evaluation view.

The listing reflects ZenMux's ongoing effort to make model evaluation more transparent for teams that compare language models across providers before deciding how to route production or experimental AI requests. Rather than publishing a static description, ZenMux pairs the model page with live platform indicators so technical teams can assess operational conditions before testing the model in their own environments.

Z.AI GLM 5.3 Flash is described as a native multimodal model from Z.ai suited for efficient coding and long-horizon agent tasks. According to the model listing, its hybrid sparse and linear attention architecture is

designed to maintain accurate long-context behavior while reducing compute overhead, a combination ZenMux said is relevant to teams running high-volume or latency-sensitive agent workflows. The current listing identifies BigModel as a provider for GLM 5.3 Flash and displays live indicators including cache hit rate, context length, latency, throughput, and recent uptime.

“Developers evaluating a new model rarely stop at capability alone,” said a ZenMux spokesperson. “With GLM 5.3 Flash, teams want to see how the model behaves under real provider conditions, including latency, throughput, cache performance, and uptime, before they commit to testing it against production prompts. ZenMux GLM 5.3 Flash is set up to put that operational context next to the model description, so evaluation can start with a fuller picture rather than a name and a benchmark score.”

ZenMux said the addition is particularly relevant for organizations where model choice shifts as projects move from prototype to production. A team may test one model for coding assistance, another for summarization, and another for agentic task execution, all while needing a consistent way to compare provider routing and performance signals. ZenMux positions its platform as a way to manage that comparison process through unified access, provider routing, and model catalog visibility, rather than requiring teams to check separate provider dashboards and documentation sets.

The GLM 5.3 Flash listing also appears alongside related Z.ai models, including GLM 5.3, GLM 5.2, and GLM 5.1, giving users a way to understand the Flash version within a broader model family while still reviewing the specific provider and performance indicators tied to it. ZenMux’s model routing environment supports access to multiple AI models through compatible API pathways and a shared interface for discovery and monitoring, a structure the company said can reduce integration work when teams compare models from different providers and helps organizations maintain a clearer record of which models are being tested and how availability changes over time.

ZenMux noted that teams should continue to perform their own quality, security, and compliance reviews before using any model in production. Model catalog data and provider indicators can support evaluation, but they do not replace internal testing against real prompts, user requirements, governance policies, and application constraints. Performance expectations can vary by workload, prompt structure, provider, and context length, which is why ZenMux surfaces operational data directly alongside the model description rather than relying solely on static specifications.

For more information about GLM 5.3 Flash routing availability, visit <https://zenmux.ai/z-ai/glm-5.3-flash>

About ZenMux:

ZenMux is an enterprise-grade large model aggregation platform with an insurance payout mechanism. The platform provides one-stop access to the latest models across providers. When issues such as poor output quality or excessive latency occur during use, our intelligent insurance detection and payout mechanism automatically compensates, addressing enterprise concerns around AI hallucinations and unstable quality.

###

For more information about ZenMux, contact the company here: ZenMuxEmberember@zenmux.ai Singapore

ZenMux

ZenMux is an enterprise-grade AI model aggregation platform with a built-in insurance payout mechanism, offering one-stop access across providers and automatic compensation for quality or latency issues.

Website: <https://zenmux.ai/>

Email: ember@zenmux.ai

