

Back to Models



DeepSeek: DeepSeek V4.1 Flash

Active 50% OFF

Compare API Request Chat

deepseek/deepseek-v4.1-flash

DeepSeek V4.1 Flash is a 552B-parameter MoE model built on a new causal encoder-decoder architecture, designed for higher capability, faster reasoning, higher throughput, and lower serving cost. It natively supports multimodal visual understanding and delivers flagship-level intelligence with significantly reduced KV cache requirements, making it well suited for high-throughput and cost-sensitive agentic workloads.

By : DeepSeek Input Type : Output Type : Publish time : 2026-09-10

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
Alibaba Cloud alibaba	0.97%	¥0.075-¥0.15 / M tokens	¥0.3-¥0.6 / M tokens	Read: ¥0.015-¥0.03 0.0075-0.015 / M tokens Write: - / M tokens	1M	4.45s	60.7tps
DeepSeek deepseek	98.5%	¥0.075-¥0.15 / M tokens	¥0.3-¥0.6 / M tokens	Read: ¥0.003-¥0.006 0.0015-0.003 / M tokens Write: - / M tokens	1M	1.14s	141tps
Baidu baidu	-	¥0.075-¥0.15 / M tokens	¥0.3-¥0.6 / M tokens	Read: ¥0.003-¥0.006 0.0015-0.003 / M tokens Write: - / M tokens	1M	4.8s	102tps

ZenMux Launches DeepSeek V4.1 Flash for High-Throughput Reasoning

September 22, 2026

Singapore - September 22, 2026 -

ZenMux has added DeepSeek V4.1 Flash to its routed model catalog, with provider options listed for DeepSeek, Alibaba Cloud, and Baidu. The model page records a publish date of September 10, 2026 and identifies a one-million-token context window. The launch gives developers a new option for testing high-throughput reasoning, multimodal understanding, and agent-oriented workloads through the ZenMux platform.

According to the ZenMux model page, DeepSeek V4.1 Flash is a 552-billion-parameter mixture-of-experts model based on a causal encoder-decoder architecture. The listing describes the model as designed for higher throughput, faster reasoning, and lower serving cost while reducing key-value cache requirements ? a design intended to lower memory usage and serving cost during inference. These statements reflect the specifications published on ZenMux and should be assessed by development teams against their own workloads rather than treated as universal performance results. The mixture-of-experts structure allows the model to activate only a subset of parameters for a given input, an approach commonly used to balance large

total parameter counts against practical inference costs. This architectural choice is consistent with the page's emphasis on throughput and reduced serving overhead relative to comparably sized dense models.

The ZenMux DeepSeek V4.1 Flash page also describes native visual understanding. That capability gives application teams an additional option when a workflow combines written instructions with image-based material. Practical performance can vary by task, prompt design, provider implementation, and input quality, so production use should include representative testing, error analysis, and human review where appropriate.

ZenMux lists three routes for the model: DeepSeek, Alibaba Cloud, and Baidu. Multi-provider availability can be useful for organizations that need to compare regional access, capacity, response behavior, or operational resilience. This makes DeepSeek V4.1 Flash one of the more route-diverse models currently listed on ZenMux, offering additional flexibility for teams with regional access or supply-chain redundancy requirements.

The one-million-token context specification supports experiments involving large collections of text, extended agent histories, repository-scale coding tasks, or research that draws from many source documents. Large context windows do not eliminate the need for careful information selection. Sending more material can increase processing time and cost, and irrelevant context can reduce output quality. Developers can combine context management with retrieval, summarization, caching, and request-level controls based on the needs of the application.

ZenMux publishes route-level information on the dedicated page, including current price fields, cache data, response measurements, and throughput observations. This announcement therefore focuses on the model's availability, listed architecture, context capacity, input capability, and provider routes instead of presenting a temporary metric as a guaranteed result.

Developers can review ZenMux DeepSeek V4.1 Flash for the current model identifier and route details. The page serves as the source for updated provider availability and platform measurements, while the application team determines whether the model meets its quality, reliability, cost, and compliance requirements.

The addition continues ZenMux's expansion of a catalog built around access to models from multiple providers. By listing routes and technical specifications in a consistent format, the platform aims to make model comparison and integration planning more manageable for developers.

For more information, visit <https://zenmux.ai/deepseek/deepseek-v4.1-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

AI

Simplified & Assured.

ZenMux