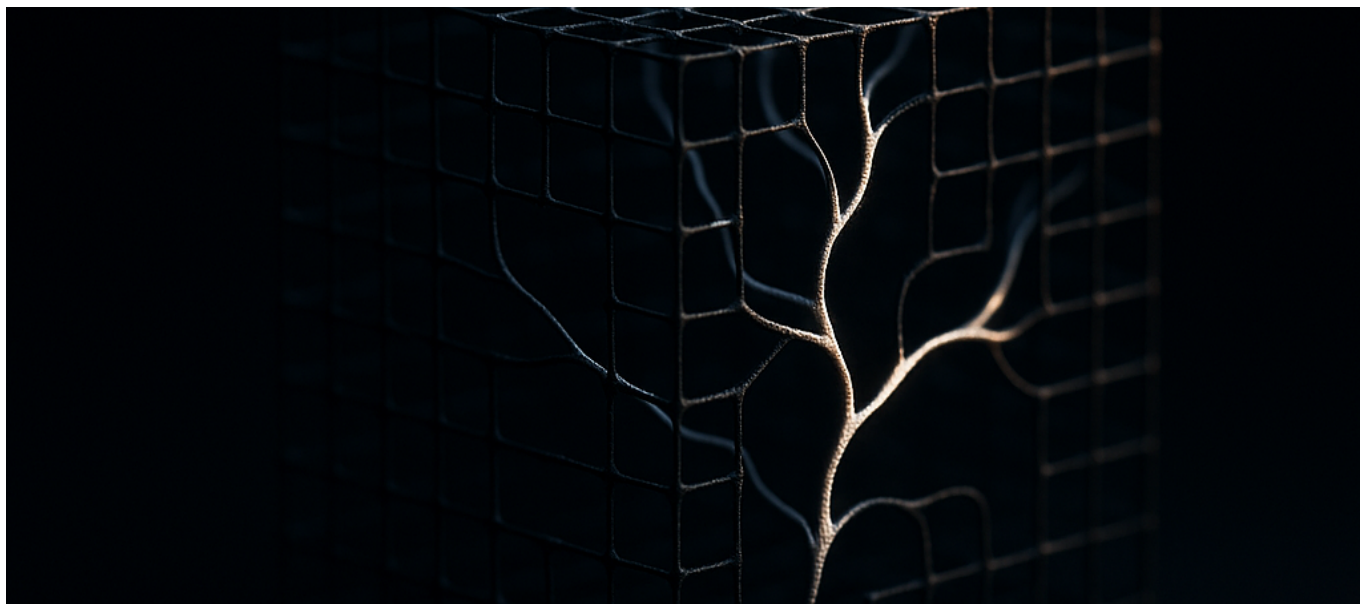




NVIDIA Nemotron 3.5 Lightning Launches at 1,200
Tokens/Second—30B MoE Outpaces Gemma 4 by 29× and
Qwen 3.6 by 35% on Agent Tasks



NVIDIA Nemotron 3.5 Lightning Launches at 1,200 Tokens/Second—30B MoE Outpaces Gemma 4 by 29× and Qwen 3.6 by 35% on Agent Tasks

NVIDIA just shipped a 30B parameter model that outputs 1,200 tokens per second—29 times faster than Gemma 4 26B on identical hardware. This isn't a chatbot play; it's infrastructure for AI agents that never sleep.

The News: NVIDIA Enters the Agent Wars

On August 11, 2026, NVIDIA released [Nemotron 3.5 Lightning](#), positioning it as the first open-source model purpose-built for always-on AI agents. The model ships with 30 billion total parameters but activates only 3 billion during inference—a mixture-of-experts (MoE) architecture that trades parameter bloat for raw speed.



NVIDIA Nemotron 3.5 Lightning Launches at 1,200 Tokens/Second—30B MoE Outpaces Gemma 4 by 29× and Qwen 3.6 by 35% on Agent Tasks

The performance numbers demand attention. According to [Artificial Analysis benchmarks](#), the model delivers 1,200.57 tokens per second at p50 latency. For comparison, Gemma 4 26B manages 41.45 tokens per second on the same prompt. That's not a 29% improvement; it's a 29× improvement.

NVIDIA licensed the model under OpenMDW-1.1, permitting unrestricted commercial use. The [model card](#) confirms a 1 million token context window and single-GPU deployment. The company also released NeMo Switchyard, an open-source routing library designed for orchestrating multiple models in agentic pipelines.

The model achieves 86% accuracy on PinchBench, the increasingly standard benchmark for measuring agent productivity in multi-step task completion. Against Qwen 3.6 35B, Nemotron 3.5 Lightning completed 10,000 agent tasks 35% faster while maintaining comparable accuracy. On the AA-Omniscience Non-Hallucination metric, it scores 69.9%—the next best model in its class sits at 50.3%.

Why This Matters: The Economics of Always-On Agents

The AI industry spent 2024 and 2025 obsessing over model intelligence. The 2026 battleground has shifted to model efficiency—specifically, the cost of keeping agents running 24/7.

Consider the math. A customer service agent handling 1,000 conversations daily generates roughly 2 million output tokens. At typical cloud inference rates of \$0.03 per 1,000 tokens, that's \$60 per day per agent, or \$21,900 annually. Multiply by 100 agents and you're looking at \$2.19 million in inference costs alone.

Speed directly impacts this equation. A model that outputs tokens 4× faster can process 4× the workload on identical infrastructure. NVIDIA's claim of 1,200 tokens/second means a single GPU running Nemotron 3.5 Lightning can handle the throughput that previously required 29 GPUs running Gemma 4. The infrastructure savings compound aggressively.

The model that wins the agent market won't be the smartest—it'll be the one enterprises can afford to leave running indefinitely.



NVIDIA Nemotron 3.5 Lightning Launches at 1,200 Tokens/Second—30B MoE Outpaces Gemma 4 by 29× and Qwen 3.6 by 35% on Agent Tasks

This release represents NVIDIA's first serious move into the open-weights foundation model space. As [CNBC reported](#), it's also the company's first open-source AI model—a strategic pivot from a hardware vendor that previously let others build the software stack.

The timing matters. Google's Gemma 4 and Alibaba's Qwen 3.6 have dominated the open-source agent model conversation for the past year. NVIDIA's entry directly targets their installed base with a model optimized for NVIDIA hardware—creating a vertical integration play that neither Google nor Alibaba can replicate.

Winners and Losers

Winners:

- Enterprises running NVIDIA infrastructure who gain a purpose-built agent model without licensing friction
- Startups building always-on agent products who can now project significantly lower unit economics
- The open-source ecosystem, which gains another high-quality model with permissive licensing

Losers:

- Inference API providers whose pricing assumed lower throughput models
- Qwen and Gemma adoption in agent-specific use cases where speed trumps marginal intelligence gains
- Companies that over-invested in context window expansion rather than throughput optimization

Technical Deep Dive: How 30B Parameters Run Like 3B

The architecture deserves scrutiny because it explains both the speed gains and the tradeoffs.

Mixture-of-experts models partition their parameters into specialized subnetworks. During inference, a routing mechanism activates only the relevant experts for each token. Nemotron 3.5 Lightning implements this with a 30B total parameter budget



NVIDIA Nemotron 3.5 Lightning Launches at 1,200 Tokens/Second—30B MoE Outpaces Gemma 4 by 29× and Qwen 3.6 by 35% on Agent Tasks

but only 3B active parameters per forward pass.

This design creates a favorable arithmetic intensity profile. Memory bandwidth—not compute—typically bottlenecks transformer inference at smaller batch sizes. By reducing active parameters 10×, the model demands proportionally less memory movement while maintaining access to the full parameter space when needed.

The 1 million token context window adds another dimension. Standard attention mechanisms scale quadratically with sequence length, but NVIDIA’s implementation clearly uses sparse or linear attention variants to make megatoken contexts practical on single-GPU deployment. The [model card](#) confirms single-GPU operation but doesn’t detail the specific attention architecture—an implementation detail worth investigating before production deployment.

Benchmark Breakdown

Let’s contextualize the published numbers:

PinchBench (86% accuracy): This benchmark evaluates multi-step task completion—the bread and butter of agentic workloads. An agent that scores 86% completes the requested action correctly 86% of the time without human intervention. In production, this translates to 14% of tasks requiring fallback or escalation.

AA-Omniscience Non-Hallucination (69.9%): This metric measures how often the model refuses to answer rather than hallucinating incorrect information. At 69.9% versus 50.3% for the next-best model, Nemotron 3.5 Lightning demonstrates markedly better calibration. For agents making autonomous decisions, hallucination rates determine liability exposure.

MMLU Pro (81.94%): The standard knowledge benchmark confirms the model hasn’t sacrificed general capability for speed. For context, GPT-4’s initial MMLU Pro score was 86.4%. The 4.5-point gap represents acceptable degradation for most production use cases.

Intelligence Index (24 vs. 15): Artificial Analysis’s composite metric shows a 60% improvement over Nemotron 3 Nano. The distillation from Nemotron 3 Ultra evidently preserved substantial capability while dramatically reducing inference cost.



NVIDIA Nemotron 3.5 Lightning Launches at 1,200 Tokens/Second—30B MoE Outpaces Gemma 4 by 29× and Qwen 3.6 by 35% on Agent Tasks

Time to First Token (1.02s p50): This latency matters for interactive applications. Sub-second TTFT enables responsive agent UX. The 1.11s p95 suggests consistent performance without long-tail latency spikes that degrade user experience.

The NeMo Switchyard Component

NVIDIA's simultaneous release of NeMo Switchyard deserves attention. This open-source routing library enables dynamic model selection within agentic pipelines—routing simple queries to smaller models and complex reasoning tasks to larger ones.

The strategic intent is clear: NVIDIA wants to own the orchestration layer for multi-model agent systems. Switchyard creates switching costs. Once teams build pipelines around NVIDIA's routing abstractions, migrating to alternative infrastructure becomes significantly more expensive.

For practitioners, Switchyard offers genuine utility. Intelligent routing can reduce average inference cost 50-70% by matching query complexity to model capability. But be aware that you're adopting NVIDIA's opinions about model orchestration alongside the tooling.

The Contrarian Take: What the Coverage Gets Wrong

Most coverage has framed Nemotron 3.5 Lightning as NVIDIA's entry into the foundation model wars. That's technically accurate but strategically incomplete. This release is primarily a hardware lock-in mechanism masquerading as an open-source contribution.

The model runs fastest on NVIDIA hardware by design. While the weights are open, the inference stack is optimized for CUDA. Running Nemotron 3.5 Lightning on AMD or Intel accelerators will sacrifice much of the 29× speed advantage. The permissive license invites adoption; the performance characteristics ensure that adoption happens on NVIDIA GPUs.

This isn't criticism—it's clarification. Understanding NVIDIA's incentive structure helps predict the model's evolution. Future versions will likely continue optimizing



NVIDIA Nemotron 3.5 Lightning Launches at 1,200 Tokens/Second—30B MoE Outpaces Gemma 4 by 29× and Qwen 3.6 by 35% on Agent Tasks

for NVIDIA-specific features like Transformer Engine and FP8 inference, creating growing performance gaps on competing hardware.

The 1 million token context window is both overhyped and underhyped.

Overhyped because most agent tasks don't require megatoken contexts. The average agentic workflow involves dozens of tool calls with modest context, not novel-length document processing. Underhyped because the few workloads that do require extreme context—codebase analysis, legal document review, research synthesis—have been locked out of efficient inference until now.

The 35% speed improvement over Qwen 3.6 matters more than the 29× improvement over Gemma 4. The Gemma comparison grabs headlines but compares models with different optimization profiles. The Qwen comparison reflects realistic competitive positioning: two models optimized for similar workloads, with Nemotron winning on throughput while matching accuracy.

The biggest underreported story is the hallucination metric. A 69.9% versus 50.3% gap in hallucination resistance represents a near-40% relative improvement in reliability. For autonomous agents making decisions without human review, this reliability differential changes the viability calculation for entire categories of deployment.

Practical Implications: What Should You Actually Do?

Let's get specific about how this release should influence decision-making.

If You're Evaluating Agent Infrastructure

Add Nemotron 3.5 Lightning to your benchmark suite immediately. The model's combination of speed, accuracy, and hallucination resistance makes it a legitimate contender for production agent deployments.

Run your specific workloads on your specific hardware before making decisions. Published benchmarks use controlled conditions that may not reflect your query distribution, context lengths, or concurrent user loads. The 1,200 tokens/second figure represents p50 performance on NVIDIA's test infrastructure—your mileage will vary based on batch size, prompt complexity, and GPU generation.



NVIDIA Nemotron 3.5 Lightning Launches at 1,200 Tokens/Second—30B MoE Outpaces Gemma 4 by 29× and Qwen 3.6 by 35% on Agent Tasks

Consider the TCO implications carefully. A model that runs 4× faster on identical hardware delivers 4× better unit economics—but only if that hardware is NVIDIA GPUs. If your infrastructure runs on alternative accelerators, the performance gap may narrow or reverse.

If You're Building Agent Products

The 86% PinchBench accuracy creates interesting product design decisions. With 14% of tasks requiring intervention, your UX must gracefully handle failure cases. The speed advantage means you can potentially run multiple inference passes per user interaction—implementing verification or consensus mechanisms that improve end-to-end reliability.

The hallucination resistance metric (69.9%) enables more autonomous decision-making than previously practical. Tasks that previously required human-in-the-loop verification may become fully automatable. Audit your current human review requirements and identify candidates for reduction.

NeMo Switchyard is worth evaluating if you're routing between multiple models. The abstraction overhead is non-trivial, but the ability to dynamically select models based on query characteristics can meaningfully impact both cost and quality.

If You're a Platform Team

Prepare for increased demand for NVIDIA inference capacity. Models optimized for specific hardware accelerate adoption of that hardware. Budget conversations should anticipate that teams previously running smaller models may request capacity for larger MoE architectures.

Evaluate your model serving infrastructure against MoE requirements. Mixture-of-experts models have different memory access patterns than dense transformers. Naive serving configurations may leave significant performance on the table. TensorRT-LLM and vLLM both support MoE optimization—ensure your deployments use appropriate backends.

The 1 million token context window creates capacity planning challenges. A single long-context inference can consume GPU memory that would otherwise serve multiple short-context requests. Consider implementing request routing that segregates long-context workloads onto dedicated capacity.



NVIDIA Nemotron 3.5 Lightning Launches at 1,200 Tokens/Second—30B MoE Outpaces Gemma 4 by 29× and Qwen 3.6 by 35% on Agent Tasks

Code Worth Exploring

Access the model through NVIDIA’s Build API for initial experimentation:

Start with your actual production queries, not generic benchmarks. The speed advantages are real, but task-specific accuracy requires validation on task-specific inputs.

For NeMo Switchyard evaluation, NVIDIA’s examples demonstrate basic routing patterns. The framework’s value proposition depends heavily on your specific multi-model requirements—teams running single-model stacks gain little from routing abstraction.

Forward Look: The Next 12 Months

This release signals several trends that will compound over the coming year.

Agent-optimized models become a distinct category. The industry has discussed “reasoning models” and “coding models” as specializations. Nemotron 3.5 Lightning establishes “agent models” as a recognized optimization target with specific benchmarks (PinchBench), specific metrics (hallucination resistance), and specific architectural choices (MoE for throughput, extended context for multi-turn).

Expect Google and Alibaba to release explicit agent-optimized variants within 90 days. Neither company will cede this market positioning without response. The competitive pressure should accelerate open-source agent model quality across all providers.

The MoE architecture becomes default for production inference. Dense models offer simpler training dynamics, but MoE’s inference efficiency advantages are now impossible to ignore. New model releases will increasingly adopt sparse architectures, with dense models relegated to research contexts where training simplicity matters more than deployment cost.

Hardware lock-in through software intensifies. NVIDIA demonstrated that open-source models can still create switching costs. Expect AMD and Intel to respond with their own model releases optimized for their accelerators. The “best available model” will increasingly depend on your hardware vendor, fragmenting the previously unified open-source model ecosystem.



NVIDIA Nemotron 3.5 Lightning Launches at 1,200 Tokens/Second—30B MoE Outpaces Gemma 4 by 29× and Qwen 3.6 by 35% on Agent Tasks

Always-on agents reach enterprise production. The economic barriers to continuous agent operation have constrained deployment to high-value use cases. With 4× efficiency improvements, the viable deployment surface expands dramatically. Customer service, IT helpdesk, sales qualification, and compliance monitoring agents will move from pilot to production at organizations with tolerance for the remaining 14% failure rate.

Model routing becomes a required competency. Multi-model architectures that route queries to appropriate models based on complexity, latency requirements, and cost constraints will become standard practice. Teams that treat model selection as a static configuration decision will find themselves at competitive disadvantage to organizations that dynamically optimize across model portfolios.

What to Watch

- **Nemotron 3.5 Ultra release timing:** NVIDIA described Lightning as “the first model in the Nemotron 3.5 family.” Larger variants optimized for reasoning-heavy agent tasks should follow within months.
- **PinchBench adoption as industry standard:** Benchmark adoption drives optimization. If PinchBench becomes the accepted agent benchmark, model development will optimize specifically for its evaluation criteria.
- **AMD response timeline:** AMD’s ROCm ecosystem has closed significant gaps with CUDA. An agent-optimized model targeting AMD hardware would materially change competitive dynamics.
- **Enterprise production deployments:** Pilot announcements mean little. Watch for companies reporting Nemotron 3.5 Lightning in production for revenue-generating workloads at scale.

The Broader Context: NVIDIA’s Strategic Evolution

This release marks a strategic pivot worth understanding. NVIDIA built a \$3 trillion market cap selling hardware and letting others build the software stack. Releasing production-quality open-source models changes that dynamic fundamentally.

The logic is straightforward. As inference costs decline and model quality converges, hardware margins face pressure. Vertical integration—where NVIDIA



NVIDIA Nemotron 3.5 Lightning Launches at 1,200 Tokens/Second—30B MoE Outpaces Gemma 4 by 29× and Qwen 3.6 by 35% on Agent Tasks

models run best on NVIDIA hardware—protects those margins by creating performance-based switching costs that pure commodity hardware cannot match.

This strategy creates genuine value for NVIDIA customers. Tight hardware-software integration enables optimizations impossible in generic inference frameworks. The 29× speed advantage over Gemma 4 exists precisely because NVIDIA controls both the model architecture and the execution environment.

But it also creates lock-in that careful practitioners should recognize. Betting on Nemotron 3.5 Lightning is betting on the NVIDIA ecosystem. That bet has historically paid off, but it's a bet nonetheless.

The company that makes the silicon now makes the models optimized for that silicon. The separation of concerns that defined the ML stack is collapsing.

What This Means for Your 2027 Planning

Several planning assumptions require revision:

Model inference costs will decline faster than projected. MoE architectures plus hardware-specific optimization create compounding efficiency gains. Budget forecasts assuming gradual inference cost reduction should anticipate step-function improvements.

Agent deployment timelines should accelerate. The combination of higher accuracy, better hallucination resistance, and improved economics removes several deployment blockers simultaneously. Teams that planned agent initiatives for late 2027 should evaluate moving timelines forward.

Multi-vendor model strategies face new complexity. The assumption that open-source models run equivalently across hardware vendors no longer holds. Platform teams must either standardize on a hardware vendor or maintain parallel model deployments optimized for different accelerators.

Build versus buy decisions shift toward build. Permissive licensing and single-GPU deployment make self-hosting increasingly viable for organizations with



NVIDIA Nemotron 3.5 Lightning Launches at 1,200 Tokens/Second—30B MoE Outpaces Gemma 4 by 29× and Qwen 3.6 by 35% on Agent Tasks

infrastructure competency. The case for third-party inference APIs weakens when equivalent models run efficiently on owned hardware.

Model evaluation becomes continuous, not periodic. The competitive pace of model releases—NVIDIA today, Google and Alibaba responding within weeks—requires ongoing evaluation capacity. Static model selection based on point-in-time benchmarks creates technical debt as better options emerge.

Final Assessment

Nemotron 3.5 Lightning represents NVIDIA's serious entry into production foundation models, optimized specifically for the agentic workloads that will define the next phase of enterprise AI deployment.

The numbers are real: 29× throughput advantage over Gemma 4, 35% faster task completion than Qwen 3.6, 86% agent task accuracy, and meaningfully better hallucination resistance than any model in its class. These aren't marginal improvements; they represent capability thresholds that enable new categories of deployment.

The strategic implications are equally significant. NVIDIA's vertical integration—model plus hardware plus orchestration tooling—creates competitive dynamics that pure-play model vendors and pure-play hardware vendors cannot match. The open-source licensing invites adoption while the optimization profile ensures that adoption happens on NVIDIA terms.

For practitioners, the path forward is clear: benchmark the model against your specific workloads, evaluate the infrastructure implications of adoption, and plan for a competitive landscape where agent-optimized models become the standard expectation rather than a specialized niche.

The model that runs fastest while staying accurate enough wins the agent market—and NVIDIA just made the most aggressive bid yet for that position.