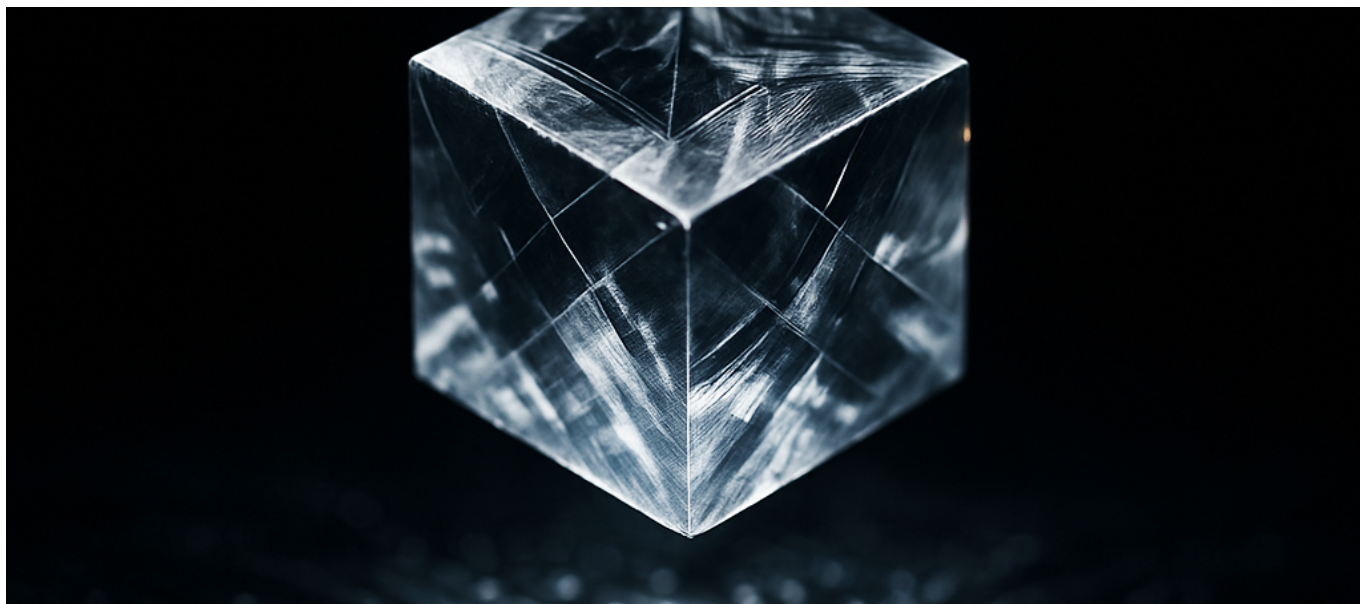




Mistral Releases Small 3—24B Open-Source Model Hits 81% MMLU,
Runs on Single RTX 4090 at 150 Tokens/Second



Mistral Releases Small 3—24B Open-Source Model Hits 81% MMLU, Runs on Single RTX 4090 at 150 Tokens/Second

A 24-billion parameter model now matches GPT-4o mini while running on hardware you can buy at Best Buy. Mistral just made the \$400 billion cloud inference market look very different.

The News: Mistral Drops a Bomb on January 31

Mistral AI [released Mistral Small 3](#) on January 31, 2025, under the Apache 2.0 license—the most permissive open-source license available. The model packs 24 billion parameters, achieves 81% accuracy on the MMLU benchmark, and runs at 150 tokens per second on a single NVIDIA RTX 4090.

The numbers demand attention. On the ARC-C reasoning benchmark, Mistral Small 3



Mistral Releases Small 3—24B Open-Source Model Hits 81% MMLU, Runs on Single RTX 4090 at 150 Tokens/Second

scores 91.3%. On GSM8K mathematical reasoning, it hits 80.7%. On HumanEval coding tasks, it reaches 84.8%. These aren't incremental improvements over previous open-source models—they represent parity with proprietary systems that cost 10x more to operate.

For those who prefer cloud deployment, Mistral offers API access at \$0.10 per million input tokens and \$0.30 per million output tokens. But the real story isn't the API pricing. The real story is that you don't need the API at all.

According to [Artificial Analysis's benchmarking](#), Mistral Small 3 performs competitively with Llama 3.3 70B and Qwen 32B—models nearly three times its size—while running at triple their inference speed. When quantized to 4-bit precision, it runs on a MacBook with 32GB of RAM. The performance ceiling for local AI inference just moved significantly upward.

Why This Release Changes the Economics of AI Deployment

The enterprise AI landscape has operated on an implicit assumption: serious AI workloads require cloud infrastructure. Mistral Small 3 challenges that assumption directly.

The Cost Calculus Shifts

Consider the economics of a typical enterprise deployment. Running GPT-4o mini through OpenAI's API at scale costs roughly \$0.15 per million input tokens and \$0.60 per million output tokens. For a company processing 10 million tokens daily, that's approximately \$2,250 per month in API costs alone—before accounting for latency, data transfer fees, and the operational overhead of managing API dependencies.

A single RTX 4090 costs \$1,599 at retail. Running Mistral Small 3 locally, that same 10 million daily tokens costs electricity—roughly \$50 per month at average US commercial rates. The hardware pays for itself in three weeks.

But the cost savings, while substantial, aren't the most important factor. The more consequential shift involves data sovereignty and operational independence.



Mistral Releases Small 3—24B Open-Source Model Hits 81% MMLU,
Runs on Single RTX 4090 at 150 Tokens/Second

Who Wins and Who Loses

Winners: Regulated Industries. Financial services firms testing fraud detection and healthcare organizations building patient triage systems have struggled with a fundamental tension. They need AI capabilities, but they operate under regulatory frameworks that make sending sensitive data to third-party APIs legally problematic or outright prohibited. HIPAA, GDPR, SOX, and PCI-DSS all create friction around cloud AI adoption.

Mistral Small 3 eliminates that friction. A healthcare provider can run diagnostic support AI entirely on-premises. A bank can deploy fraud detection without ever transmitting transaction data externally. The model doesn't make these use cases possible—they were always technically possible with open weights—but it makes them practical at a quality level that satisfies production requirements.

Winners: Edge Computing Applications. Autonomous vehicles, industrial robotics, and remote operations have struggled with the latency requirements of cloud inference. A 200ms round-trip to a cloud API is unacceptable when you're making real-time decisions. At 150 tokens per second on local hardware, Mistral Small 3 enables AI-powered decision-making in latency-critical contexts that were previously impractical.

Losers: Cloud Inference Middlemen. The venture capital thesis behind dozens of inference-as-a-service startups assumed that running AI models at scale would remain sufficiently complex that businesses would pay for managed services. When a 24B model runs on consumer hardware, that thesis weakens. Companies with technical teams capable of basic infrastructure management now have a credible alternative to inference APIs.

Losers: Proprietary Model Lock-in Strategies. OpenAI, Anthropic, and Google have all built business models around proprietary model access. Their value proposition requires that open-source alternatives remain materially inferior. Mistral Small 3 demonstrates that the gap between proprietary and open-source performance is narrowing faster than many projected.

Technical Deep Dive: How 24B Parameters Compete with



70B

The headline numbers are impressive, but the engineering decisions that enabled them deserve examination. Mistral Small 3 represents a different philosophy than simply scaling parameter counts.

Architecture Efficiency Over Raw Scale

According to [LLM Stats data](#), Mistral Small 3 achieves its benchmark performance through architectural refinements rather than brute-force scaling. The model uses a mixture of techniques that prioritize inference efficiency.

The key insight: parameter count is a poor proxy for model capability. A well-architected 24B model can match a less efficient 70B model because what matters is how effectively parameters are utilized, not how many exist. Mistral's engineering team clearly prioritized inference latency and memory efficiency as first-class design constraints.

The 150 tokens/second throughput on an RTX 4090 is particularly notable. For comparison, Llama 3.3 70B on the same hardware typically achieves 40–60 tokens per second even with aggressive quantization. Mistral Small 3 delivers roughly 3x the throughput at comparable quality. That's not a marginal improvement—it's the difference between practical local deployment and frustrating local deployment.

Benchmark Analysis: Where Mistral Small 3 Excels

The benchmark distribution reveals where Mistral focused optimization efforts:

- **MMLU (81%):** Broad knowledge reasoning. This is the headline benchmark because it correlates well with general-purpose assistant capabilities. 81% is competitive with GPT-4o mini and exceeds most open-source models in this parameter range.
- **ARC-C (91.3%):** Grade-school science questions requiring genuine reasoning rather than pattern matching. The high score here suggests robust logical reasoning capabilities.
- **GSM8K (80.7%):** Grade-school math word problems. Math reasoning has historically been a weakness of smaller models, making this score particularly significant.



Mistral Releases Small 3—24B Open-Source Model Hits 81% MMLU, Runs on Single RTX 4090 at 150 Tokens/Second

- **HumanEval (84.8%):** Code generation accuracy. For engineering teams considering local deployment for development tools, this benchmark matters more than MMLU.
- **MATH (70.6%):** Competition-level mathematics. This is notably lower than the other benchmarks, suggesting that complex mathematical reasoning remains a weak point—consistent with most models in this class.

The [n8n benchmark results](#) confirm strong performance on practical agentic tasks, which matters for teams building AI-powered workflows and automation.

The Quantization Story

Running the full 24B parameter model at full precision requires approximately 48GB of VRAM—more than any consumer GPU offers. The practical deployment story depends entirely on quantization.

At 8-bit quantization, the model fits on an RTX 4090's 24GB VRAM with room for context. At 4-bit quantization, it runs on an RTX 3090 or even a MacBook Pro with 32GB unified memory. The performance degradation from quantization is measurable but modest—typically 2–4 percentage points on benchmarks, which still leaves the model competitive with much larger alternatives.

The implication for deployment teams: quantization strategy is now a core architectural decision. The same model can serve vastly different hardware profiles depending on precision choices. A company deploying across heterogeneous infrastructure—cloud for peak capacity, edge devices for routine operations—can use identical model weights with different quantization levels.

What Most Coverage Gets Wrong

The coverage of Mistral Small 3 has focused on benchmark comparisons and hardware requirements. That framing misses the more significant implications.

This Isn't Just About Cost Savings

The “run AI locally and save money” narrative is true but incomplete. Cost savings assume



Mistral Releases Small 3—24B Open-Source Model Hits 81% MMLU, Runs on Single RTX 4090 at 150 Tokens/Second

that cloud and local deployments are functionally equivalent, with the only difference being operational expense. They aren't equivalent.

Local deployment enables use cases that cloud deployment simply cannot serve. A medical device running diagnostic AI cannot have a cloud dependency—the FDA will reject it. An automotive system making real-time decisions cannot tolerate API latency—physics doesn't wait for network round trips. A defense contractor building intelligence analysis tools cannot transmit classified data to commercial cloud providers—security clearance requirements prohibit it.

Mistral Small 3's significance isn't that it's cheaper than cloud APIs. It's that it extends production-quality AI capabilities to contexts where cloud APIs were never an option.

The Apache 2.0 License Matters More Than the Benchmarks

Mistral released Small 3 under Apache 2.0—a license that permits commercial use, modification, and redistribution with essentially no restrictions beyond attribution. This is unusual for a model at this performance level.

The licensing choice enables derivative products in ways that more restrictive licenses prevent. A company can fine-tune Mistral Small 3 on proprietary data, deploy it in commercial products, and never share the modifications. They can build entire product lines on the foundation without licensing negotiations, revenue sharing, or usage restrictions.

Compare this to models with non-commercial clauses, restrictive acceptable use policies, or “open weights but closed training” approaches. The Apache 2.0 license isn't just legally convenient—it's foundational infrastructure for building commercial AI products without platform dependency.

The “GPT-4o Mini Equivalent” Framing Undersells It

Describing Mistral Small 3 as “matching GPT-4o mini” frames it as catching up to existing capabilities. The more accurate framing: Mistral Small 3 delivers GPT-4o mini-level capabilities with deployment flexibility that GPT-4o mini fundamentally cannot offer.



Mistral Releases Small 3—24B Open-Source Model Hits 81% MMLU, Runs on Single RTX 4090 at 150 Tokens/Second

OpenAI cannot release GPT-4o mini weights for local deployment—their business model depends on API access. The comparison isn't model-to-model; it's deployment paradigm to deployment paradigm. A GPT-4o mini-equivalent model that runs locally isn't just matching OpenAI's capability—it's offering something OpenAI's capability cannot include.

Practical Implementation: What to Actually Do

For technical leaders evaluating Mistral Small 3, here's a concrete implementation framework.

Evaluation Criteria

Before deploying, assess your use case against these criteria:

Latency requirements: If your application requires sub-100ms response times, local deployment is likely superior to cloud APIs even with fast network connections. The 150 tokens/second throughput translates to roughly 6-7ms per token—fast enough for real-time applications.

Data sensitivity: If your data falls under HIPAA, GDPR, PCI-DSS, or similar regulatory frameworks, local deployment eliminates compliance complexity around third-party data processing. The model weights never leave your infrastructure, and neither does your data.

Scale variability: If your workload is highly variable—peaks and valleys in demand—cloud APIs offer elasticity that fixed hardware cannot match. If your workload is consistent, local deployment offers predictable costs regardless of volume.

Task complexity: For tasks where Mistral Small 3's benchmarks exceed your requirements—most code generation, summarization, classification, and extraction tasks—local deployment is straightforwardly superior. For tasks requiring frontier model capabilities—complex reasoning chains, novel problem-solving, or tasks where you're already pushing GPT-4's limits—you may need to supplement with cloud API calls.

Hardware Selection

For teams planning hardware procurement:



Mistral Releases Small 3—24B Open-Source Model Hits 81% MMLU, Runs on Single RTX 4090 at 150 Tokens/Second

- **Development/Testing:** Any modern GPU with 8GB+ VRAM runs the 4-bit quantized model adequately. RTX 3060 or equivalent is sufficient.
- **Production (Single Node):** RTX 4090 is the optimal cost/performance choice. 24GB VRAM accommodates 8-bit quantization with full context window. Expected throughput: 150 tokens/second.
- **Production (Scale):** For higher throughput, multiple RTX 4090s or A100 GPUs. Consider inference-optimized instances from cloud providers for burst capacity.
- **Edge Deployment:** MacBook Pro with 32GB+ unified memory runs the 4-bit quantized model. NVIDIA Jetson AGX Orin for embedded applications.

Code: Getting Started

For teams ready to experiment, the deployment path is straightforward. Using Ollama:

```
ollama pull mistral-small  
ollama run mistral-small
```

For production deployment with vLLM or similar inference servers, the model is available on Hugging Face. Standard transformer inference tooling works without modification.

For fine-tuning on domain-specific data, the Apache 2.0 license permits full weight modification. Tools like LoRA, QLoRA, and full fine-tuning are all permissible. The 24B parameter count makes fine-tuning feasible on a single A100 or multiple RTX 4090s—expensive but not prohibitive.

Architecture Patterns to Consider

Hybrid Local/Cloud: Run Mistral Small 3 locally for routine operations, route complex or edge-case queries to frontier models via API. This captures 90%+ of cost savings while maintaining access to superior capabilities when needed.

Regional Inference: Deploy local instances in each geographic region where you operate. Eliminates cross-region latency, maintains data residency compliance, and provides resilience against cloud outages.



Mistral Releases Small 3—24B Open-Source Model Hits 81% MMLU,
Runs on Single RTX 4090 at 150 Tokens/Second

Tiered Processing: Use Mistral Small 3 as a first-pass filter for classification, extraction, and summarization. Only send queries requiring deeper reasoning to larger models. Most workloads are routine; reserve expensive compute for genuinely complex tasks.

The Competitive Landscape: Where This Leads

Mistral Small 3's release accelerates trends that were already emerging. Here's where those trends point over the next 6-12 months.

Immediate Effects (Next 3 Months)

Pricing pressure on inference APIs. With a credible open-source alternative at this performance level, cloud inference pricing has a ceiling. Expect OpenAI, Anthropic, and Google to announce pricing reductions for their smaller models within weeks, not months.

Enterprise pilot programs proliferate. Companies that previously dismissed local inference as impractical will launch evaluation programs. The combination of benchmark performance and deployment flexibility makes "try it on real workloads" a low-risk proposition.

Startup strategy shifts. Inference-as-a-service startups face an existential question: if customers can run competitive models locally, what's the value of a managed inference layer? The ones that survive will differentiate on capabilities beyond raw inference—fine-tuning services, observability, compliance tooling.

Medium-Term Effects (6-12 Months)

Corporate AI infrastructure buildout. Large enterprises with consistent AI workloads will invest in dedicated inference hardware. The economics favor building over buying for organizations processing billions of tokens monthly. AWS, Azure, and GCP will see AI API revenue pressured even as their infrastructure services grow from on-premises AI deployments.

Edge AI capabilities expand dramatically. With 24B parameters running on consumer hardware, the definition of "edge" expands. Retail stores, manufacturing facilities, and



Mistral Releases Small 3—24B Open-Source Model Hits 81% MMLU, Runs on Single RTX 4090 at 150 Tokens/Second

logistics operations can deploy sophisticated AI locally. Applications that required cloud connectivity become possible in environments with intermittent or nonexistent network access.

Model efficiency becomes a competitive vector. The lesson of Mistral Small 3 is that architecture efficiency matters more than parameter count. Research investment will shift toward making models more efficient at inference time, not just more capable through scaling. The next breakthrough model might not be the biggest—it might be the most efficient at a given capability level.

The Frontier Model Question

One question remains open: does Mistral Small 3 represent the capability ceiling for efficient local deployment, or a waypoint toward higher capabilities?

The trends suggest waypoint. Model efficiency has improved faster than hardware capabilities over the past two years. If that trend continues—and there's no technical reason to expect otherwise—models matching current GPT-4-class capabilities will run on consumer hardware within 18–24 months. The cloud inference premium depends on a capability gap that is measurably shrinking.

This doesn't mean cloud AI becomes irrelevant. The largest models, the most complex reasoning tasks, and the most demanding training workloads will continue to require cloud-scale infrastructure. But the addressable market for cloud inference shrinks as local deployment becomes feasible for an expanding range of capabilities.

Vendors to Watch

Beyond Mistral, several players are positioned to benefit from or respond to this shift:

NVIDIA: Consumer GPU sales to enterprises will increase as local inference becomes practical. The RTX 4090 is not designed for datacenter deployment—it lacks the reliability features of H100 or A100—but it's good enough for many production workloads at a fraction of the cost.

Apple: The M-series chips with unified memory are uniquely positioned for local inference



Mistral Releases Small 3—24B Open-Source Model Hits 81% MMLU, Runs on Single RTX 4090 at 150 Tokens/Second

on laptops and desktops. Expect Apple to emphasize AI capabilities in hardware marketing and potentially optimize Metal for transformer inference.

AMD: ROCm compatibility for inference workloads lags CUDA, but the economics favor competition. Enterprises buying inference hardware at scale will push AMD to close the gap.

Hugging Face: The neutral platform for model distribution benefits regardless of which models win. Their inference infrastructure business faces pressure, but their model hub and training tools become more valuable as more teams experiment with open weights.

Groq and Cerebras: Custom inference hardware startups face a clarified market. If consumer GPUs handle 24B parameter models adequately, custom silicon needs to target larger models or deliver order-of-magnitude improvements—marginal efficiency gains won't justify the ecosystem lock-in.

The Bottom Line

Mistral Small 3 is not a research curiosity or a benchmark trophy. It's production-ready AI that runs on hardware you can buy today, under a license that permits any commercial use, at a performance level that satisfies real workloads.

The strategic question for technical leaders isn't whether to adopt it—that depends on your specific requirements. The strategic question is what happens to your AI strategy when local deployment becomes the default rather than the exception.

For organizations handling sensitive data, operating in regulated industries, or building latency-critical applications, the answer is clear: local-first AI is now viable. For everyone else, the cost-benefit calculation has shifted permanently toward ownership over rental.

The cloud inference era isn't ending—but the assumption that it's the only serious option is now obsolete.



Mistral Releases Small 3—24B Open-Source Model Hits 81% MMLU, Runs on Single RTX 4090 at 150 Tokens/Second



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 [📅](#)