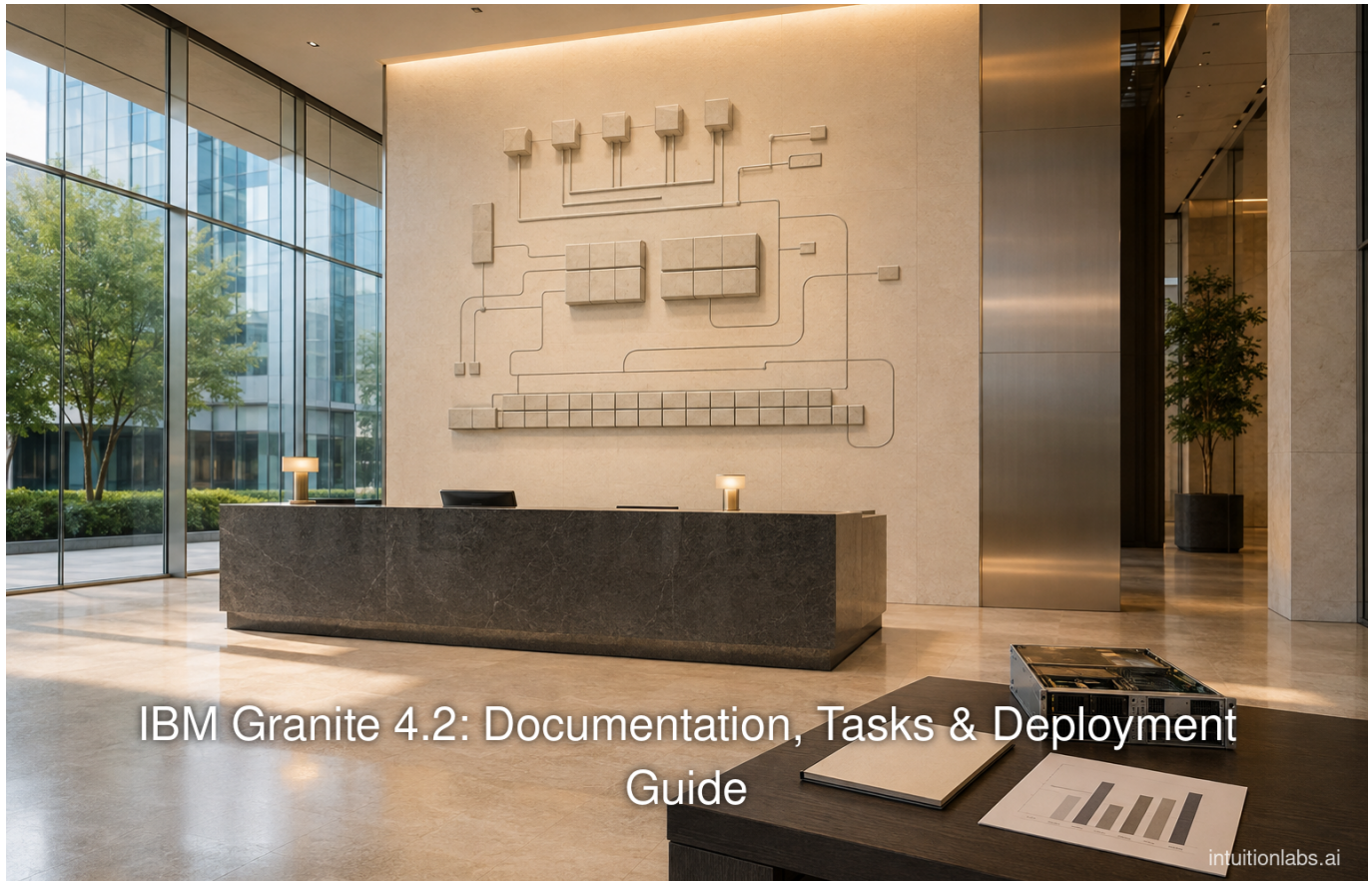


IBM Granite 4.2: Documentation, Tasks & Deployment Guide

IntuitionLabs.ai · Adrien Laurent · 2026-09-05 · [ibm granite 4.2](#) · [granite 4.2 documentation](#) · [ibm granite model](#) · [granite 4.2 deployment](#) · [granite 4.2 benchmarks](#) · [watsonx.ai](#) · [open weight llm](#) · [apache 2.0 license](#) · [enterprise ai models](#) · [granite 4.0 vs 4.2](#)



Executive Summary

IBM Granite 4.2 is a family of open-weight, dense, decoder-only reasoning language models released by IBM Research on **August 25, 2026** in three parameter sizes: **3B**, **8B**, and **30B** (see the Introduction below for the primary source). It is IBM's first Granite generation to natively combine step-by-step chain-of-thought reasoning with tool calling, positioned by IBM as "purpose-built for the agentic workflows that today's enterprise use cases require." All three checkpoints are released under the permissive **Apache License 2.0** ([Apache License 2.0](#)). Commercial use is permitted subject to the license terms, including its redistribution and notice obligations.

Architecturally, Granite 4.2 marks a deliberate departure from Granite 4.0's October 2025 hybrid Mamba-2/transformer Mixture-of-Experts (MoE) design (see the Version History section below). Granite 4.2 uses a dense, decoder-only transformer with Grouped Query Attention and Rotary Position Embeddings, continuing an architectural line IBM began with the intermediate Granite 4.1 release in April 2026 (research.ibm.com). All three sizes natively support a **128,000-token** context window, and IBM's checkpoint-specific model cards document a

long-context extension to **512K tokens** for the **3B**, **8B**, and **30B** models. On IBM's own vendor-reported benchmarks, detailed in the Data Analysis section below, Granite 4.2 30B scores **89.17** on **AIME25 math reasoning**, **61.39** on the BFCL v4 tool-calling benchmark, and **57.00** on SWE-Bench Verified, each ahead of the 8B variant.

For deployment, IBM optimizes Granite 4.2 for **vLLM version 0.20 or later** (ibm.com) and also documents SGLang, Ollama, LM Studio, Docker Model Runner, and IBM's own hosted **watsonx.ai** platform as supported paths. IBM has not published explicit minimum GPU/VRAM figures for any of the three sizes. IBM's public [watsonx.ai](https://ibm.com/watsonx/pricing) pricing page describes IBM foundation-model pricing as including pay-as-you-go pricing per million tokens and hourly rates for on-demand model hosting and deployment ([IBM watsonx.ai pricing](https://ibm.com/watsonx/pricing)). Quantized FP8, NVFP4, MXFP4, and GGUF variants are published for self-hosted, resource-constrained deployment.

This report documents Granite 4.2's release checkpoint, architecture, benchmark evidence, deployment paths, licensing terms, and competitive position against similarly sized open-weight alternatives such as Meta's Llama 4, Alibaba's Qwen3 (qwenlm.github.io), and **Mistral AI's Mistral Large 3**, all of which are also released under permissive or semi-open licenses (mistral.ai). Enterprises evaluating Granite 4.2 for regulated or agentic workflows should treat IBM's own vendor-reported benchmarks as a starting point rather than an independent verification, confirm current per-token pricing directly against the [watsonx.ai](https://ibm.com/watsonx) console before budgeting a production deployment.

Introduction and Background

IBM Research introduced **Granite 4.2** on **August 25, 2026**, describing it in its own release post as an update to the Granite language model family that adds native reasoning to models built for enterprise agents (research.ibm.com). IBM's technical writeup calls it "our first family of dense, decoder-only reasoning LLMs," available in three sizes: 3B, 8B, and 30B parameters (huggingface.co). The family is part of IBM's broader Granite line of open, enterprise-oriented foundation models, which the project's GitHub repository describes as spanning "compact edge-deployable models to large-scale reasoning systems" (github.com).

This report examines Granite 4.2 as of **September 5, 2026**, roughly two weeks after its public release, drawing on IBM's official documentation site, its Hugging Face model cards and technical blog, its GitHub repository, and independent technical press coverage. Every price, benchmark score, and availability claim below carries the date it was observed, because model documentation, pricing pages, and quantization catalogs change frequently. Readers using this report to plan a deployment should re-verify volatile figures, particularly pricing, against IBM's live [watsonx.ai](https://ibm.com/watsonx) console rather than relying on a point-in-time snapshot.

The report is organized to answer the practical questions an engineering or procurement team is likely to ask: what Granite 4.2 actually is architecturally, how it differs from the Granite 4.0 and 4.1 generations that preceded it, what infrastructure is required to run it, what it costs and how it is licensed, what IBM and independent observers say it is good for, and how its published specifications compare with other open-weight enterprise models on the market.

Product and Platform Architecture

Granite 4.2 departs from the hybrid architecture IBM used in Granite 4.0 and instead uses, in IBM's own words, "a decoder-only dense transformer architecture" (huggingface.co). The three released checkpoints, **granite-4.2-3b**, **granite-4.2-8b**, and **granite-4.2-30b**, are post-trained on top of Granite 4.1 base models. IBM's model card architecture class is listed as **GraniteForCausalLM**, described as a decoder-only dense transformer built for

reasoning. The 3B configuration uses Grouped Query Attention (GQA) with 40 attention heads and 8 key-value heads, and Rotary Position Embeddings with a theta value of 10,000,000 (specifications per the same technical blog). Precision across the family is **bfloat16**, and input/output embeddings are kept separate rather than tied.

All three sizes natively support a **128K-token** context window, and IBM's checkpoint-specific model cards document a long-context extension to **512K tokens** for the **3B**, **8B**, and **30B** models.

Each model's post-training lineage is stated explicitly on its Hugging Face card: **Granite-4.2-3B** is "post-trained from Granite-4.1-3B-Base through a rigorous multi-stage pipeline," with the base model field on the same card reading "Granite-4.1-3B-Base" (huggingface.co), and the GitHub README independently confirming that "the Granite 4.2 dense models are post-trained on top of Granite 4.1 base models" (github.com). IBM Research's own launch blog, by contrast, describes the release as "building on the Granite 4.0 foundation models" (research.ibm.com). The checkpoint-specific **model card** and GitHub README state that the Granite 4.2 dense models are post-trained on top of Granite 4.1 base models. IBM's **technical article** states that each model is pre-trained from scratch as part of its training pipeline.

Post-training combines supervised fine-tuning on roughly **100 billion tokens** across 7.2 million samples with reinforcement learning that "applies multi-phase, multi-environment reinforcement learning using Group Relative Policy Optimization" across math, code, science, instruction-following, tool use, and structured-output environments (huggingface.co). The 8B and 30B sizes receive an additional agentic-RL stage covering software engineering and terminal-based tasks that the 3B model skips. Training used 1 trillion tokens of synthetic code generated by IBM's internal "CodeAlchemy" pipeline, and data-quality filtering used GPT-OSS-120B and Gemma 4 as automated judge models (research.ibm.com). IBM trained the family on an **NVIDIA GB200 NVL72** cluster hosted by CoreWeave ([Hugging Face technical writeup](https://huggingface.co)), and the models include a speculative-decoding layer that IBM says "allows them to output text faster" and at lower serving cost (research.ibm.com).

A distinguishing feature of Granite 4.2 is its support for three selectable thinking modes, chosen via chat-template parameters: a default thinking mode with visible chain-of-thought, a non-thinking mode, and a low-effort mode. IBM documents this capability as reasoning-augmented tool calling, independently described by Unite.AI as the model reasoning "about which tool to call and why before emitting the call" (www.unite.ai) in an OpenAI-compatible function-calling format. The family supports 12 languages for multilingual dialog: "English, German, Spanish, French, Japanese, Portuguese, Arabic, Czech, Italian, Korean, Dutch, and Chinese" (huggingface.co).

Table 1 below summarizes the published architectural specifications of the three Granite 4.2 sizes, drawn from the sources cited in the preceding paragraphs.

Specification	Granite-4.2-3B	Granite-4.2-8B	Granite-4.2-30B
Architecture	Decoder-only dense transformer (reasoning)	Decoder-only dense transformer (reasoning)	Decoder-only dense transformer (reasoning)
Base checkpoint	Granite-4.1-3B-Base	Granite-4.1-8B-Base	Granite-4.1-30B-Base
Native / extended context	128K native / 512K extended	128K native / 512K extended	128K native / 512K extended
Precision	bfloat16	bfloat16	bfloat16

Specification	Granite-4.2-3B	Granite-4.2-8B	Granite-4.2-30B
Agentic-RL stage	Not applied	Applied	Applied
Positioning	Edge / resource-constrained deployment	General-purpose enterprise reasoning	Flagship reasoning and agentic workloads
License	Apache 2.0	Apache 2.0	Apache 2.0
Release date	August 25, 2026	August 25, 2026	August 25, 2026

As the table shows, the three sizes share an identical architecture family and license but differ in agentic-RL training and intended deployment footprint; the 3B model is explicitly the edge-oriented member of the family, while the 8B and 30B sizes receive the additional software-engineering and terminal-agent reinforcement learning that the smallest model skips.

Version History: Granite 4.0, 4.1, and 4.2 Compared

IBM's Granite line has moved through three distinct architectural phases in roughly eleven months. **Granite 4.0**, announced **October 2, 2025**, was described by IBM as "the next generation of IBM language models," introducing a hybrid architecture combining Mamba-2 state-space layers with "conventional transformer blocks sequentially in a 9:1 ratio" across three of its four launch models ([ibm.com](#)). **Granite-4.0-H-Small** was released as "a hybrid mixture of experts (MoE) model with 32B total parameters (9B active)" and **Granite-4.0-H-Tiny** as a hybrid MoE with 7B total parameters (1B active), while IBM also shipped a non-hybrid **Granite-4.0-Micro**, a 3B dense model with a conventional attention-driven transformer architecture, for platforms that did not yet support the hybrid design (specifications per the same announcement). IBM's launch claims for Granite 4.0 included that the smallest models "significantly outperform Granite 3.3 8B" despite being smaller, an "over 70% reduction in RAM needed to handle long inputs" versus conventional transformer models, and training on "the same carefully compiled 22T-token corpus of enterprise-focused training data" across all four models, each trained on samples up to 512K tokens of context ([ibm.com](#)). [ISO/IEC 42001:2023](#) specifies requirements for establishing, implementing, maintaining, and continually improving an artificial intelligence management system within organizations.

Granite 4.1, released **April 29, 2026**, reversed course on the hybrid design. IBM's research blog states the new dense architecture "consistently matches or outperforms the Granite 4.0 32B Mixture[-of-Experts model]" while using a smaller 8B footprint ([research.ibm.com](#)), a claim echoed in IBM's Hugging Face announcement for that release. IBM distinguished engineer Rameswar Panda is quoted explaining that the dense, non-reasoning 4.1 design "delivers competitive instruction-following and tool-calling performance without relying on long chains of thought" ([research.ibm.com](#)), framing the move away from hybrid MoE as a simplicity-for-flexibility tradeoff rather than a pure capability regression.

Granite 4.2, four months later, kept the dense architecture and added explicit reasoning on top of it. IBM's technical blog states plainly: "earlier Granite releases were strong instruction-following assistants; Granite 4.2 adds explicit reasoning" ([huggingface.co](#)), and separately notes that "the pre-training recipe closely follows the previous generation," pointing readers to the Granite 4.1 blog for pretraining detail. IBM Research frames the dense architecture's persistence into 4.2 as a portability choice: "its dense architecture supports broad compatibility, and its multiple sizes give teams" flexibility across deployment targets ([research.ibm.com](#)). Alongside the language

models, IBM released companion speech models numbered **5.0**, not 4.2, because IBM considered them a structural leap from the previous 4.1 models, at 470 million parameters, making them among the smallest speech models IBM has shipped.

IBM's repository states that the released Granite 4.2 dense models are post-trained on top of Granite 4.1 base models ([github.com](#)). No IBM source reviewed for this report publishes a direct, side-by-side benchmark table comparing Granite 4.2 scores against Granite 4.0 or 4.1 scores on identical tests; the generational comparisons above are qualitative claims made in IBM's own announcement text, not a reproducible numeric comparison, and should be read as vendor-stated positioning rather than independently verified progress.

Deployment Requirements and Infrastructure Options

IBM's [Granite documentation page](#) says Granite 4.2 is "optimized for deployment with vLLM (v0.20+)" and also warns that the documentation site is no longer updated; operational readers should confirm serving requirements in the current [Hugging Face model cards](#) and [GitHub repository](#), and the 30B model's Hugging Face card notes a custom reasoning parser "included in this repository (requires vLLM v0.20+)" together with a dedicated tool-calling parser, "qwen3_coder" ([huggingface.co](#)). For Granite 4.2 deployment, IBM's model-specific documentation says the model is optimized for vLLM (v0.20+). For high-throughput serving, IBM also documents **SGLang (v0.5.18+)** as a supported backend, pointing to an external cookbook containing a "Docker, H200/B200 launch matrix" for GPU sizing rather than publishing its own minimum VRAM figures. Across every deployment guide reviewed for this report, **no explicit minimum GPU or VRAM requirement figure** is published by IBM for any of the three sizes; this is a documented gap in IBM's own materials, not an estimate this report is willing to supply in its place.

For simpler, non-containerized use, Granite 4.2 can be pulled through **Docker Model Runner** directly (`docker model run hf.co/ibm-granite/granite-4.2-30b`), through **Ollama**, or through **LM Studio**, which IBM describes as a beginner-friendly desktop application for running large language models. IBM's Ollama guide states plainly that "larger models give better results but require more resources" and separately warns that "by default, Ollama runs models with a short context length to save memory," requiring a manual override to use Granite 4.2's 128K native context window and the documented 512K long-context extensions for the **3B**, **8B**, and **30B** checkpoints ([ibm.com](#)); the same guide's stated context ceiling, "the largest supported context for Granite 4.0 models is 131072 (128k)," still references the 4.0 generation rather than the newer 4.2 figures at the time this page was accessed. For self-hosted Transformers-based inference, the 30B model card's quick-start code loads the model with `device_map="cuda"` and `torch_dtype=torch.bfloat16` ([huggingface.co](#)), a GPU-oriented example that does not itself document a benchmarked CPU inference path.

IBM's containerized vLLM guide lists an **NVIDIA GPU with drivers installed** as a hardware prerequisite for that specific path, while its cross-partner documentation describes a much broader hardware surface for Ollama-based deployment: "GPUs from AMD, Intel, and Nvidia, ARM devices from Qualcomm" are supported through that channel ([ibm.com](#)). For organizations preferring a hosted option, IBM positions **watsonx.ai** as an option to experience the Granite dense models without self-hosting, describing it as an enterprise-grade AI studio rather than a bare model host. Additional documented distribution and hosting channels include **Docker Hub**, **Dell Pro AI Studio**, described as "a complimentary toolkit" for on-device deployment ([ibm.com](#)), and **AWS**, where customers can discover and subscribe to Granite models through Bedrock Marketplace and SageMaker. IBM Research's own launch post reiterates that the family is "built for deployment across cloud, on-premises, and edge environments" and lists platform integrations at launch including AnythingLLM, Artificial Analysis, CoreWeave Inference, DeepInfra,

GitHub, Hugging Face, LM Studio, Ollama, OpenRouter, RadixArk, Replicate, and watsonx (research.ibm.com). For memory-constrained deployment, IBM has published quantized **FP8**, **NVFP4**, and **MXFP4** variants using LLM Compressor, alongside GGUF conversions for llama.cpp-based runtimes (huggingface.co).

Licensing, Pricing, and Availability

Granite 4.2's language model weights are released under the **Apache License, Version 2.0**, dated January 2004 in its standard text (github.com). IBM's GitHub repository identifies Apache 2.0 as the distribution license. Commercial use is permitted under the Apache 2.0 grant, subject to its terms and conditions; when redistributing, recipients must receive a copy of the license and modified files must carry prominent change notices ([Apache License 2.0](https://www.apache.org/licenses/LICENSE-2.0)). As a standard Apache 2.0 grant, the license text itself confers "a perpetual, worldwide, non-exclusive, no-charge, royalty-free, irrevocable copyright license to reproduce, prepare Derivative Works" of, publicly display, and distribute the models, with or without modification (github.com). IBM's watsonx.ai documentation states: "Only IBM foundation models that you access from watsonx.ai are indemnified by IBM" ([IBM documentation](https://www.ibm.com/docs/en/watsonx/1.0.0?topic=watsonx-ai-models)). IBM's Granite product page places one explicit usage restriction on model output: users may not "engage in or facilitate any action to generate output that infringes" third-party copyright ([ibm.com](https://www.ibm.com)).

Granite 4.2 is distributed across multiple channels beyond IBM's own site. **Ollama's** official model library page for "granite4.2" confirms the models are "released under Apache 2.0 license" on that platform as well (ollama.com). On pricing, IBM's public **watsonx.ai** pricing page says that IBM foundation models include pay-as-you-go pricing per million tokens and hourly rates for on-demand model hosting and deployment ([IBM watsonx.ai pricing](https://www.ibm.com/pricing/watsonx)). Organizations evaluating watsonx.ai hosted costs should confirm current pricing directly with IBM before budgeting a production deployment.

Because the weights are Apache 2.0 licensed, organizations may download, fine-tune, and deploy them commercially subject to that license's terms; when redistributing, they must comply with its license, modification-notice, and applicable NOTICE-file obligations. The practical cost of running Granite 4.2 is therefore driven by infrastructure choice (self-hosted GPU capacity versus IBM's or a partner's hosted per-token pricing) rather than a software license fee. No source reviewed publishes a direct dollar comparison of self-hosting costs against watsonx.ai hosted costs for Granite 4.2 specifically; this report treats that comparison as an unresolved gap rather than estimating figures IBM has not published.

Use Cases and Enterprise Functional Capabilities

IBM positions Granite 4.2 as "purpose-built for the agentic workflows that today's enterprise use cases require" (research.ibm.com), with a specific emphasis on software-engineering agents that can navigate codebases, handle multi-step development tasks, and operate terminal-based tooling. Unite.AI's independent description of the underlying training methodology notes that in the software-engineering RL stage, "the model edits real repositories and passes only if the hidden test suite passes" (www.unite.ai), an outcome-based reward design rather than one graded on stylistic similarity to reference solutions. The model card for the 3B size lists its intended use cases directly as reasoning, code generation, tool calling, agentic workflows, and multilingual dialog.

Trade press coverage adds detail on the breadth of positioned use cases. MarkTechPost's launch-day summary lists IBM-stated applications spanning "terminal and DevOps automation, deep-research and search agents, long-document RAG" (www.marktechpost.com) (retrieval-augmented generation), and names target industries including "software and developer tooling, financial services, healthcare, telecom, public sector" (www.marktechpost.com). IBM's stated intent for the overall family is "helping enterprises build agents that can reason, act, and adapt during real-life workflows" (research.ibm.com), and the companion Granite Speech 5.0 models released alongside the language family are positioned for high-volume call-center transcription workloads. None of the sources reviewed for this report name a specific enterprise customer or a completed production deployment of Granite 4.2; IBM's own materials list distribution and hosting platforms rather than named end users, and this report does not supply a case study section as a result, consistent with the requirement to avoid weak or unverifiable material rather than fill a section for its own sake.

Developer-community reception, as observed roughly two weeks post-launch, is still early-stage by conventional open-source adoption metrics. The organization's discussion board included practical follow-up questions such as "are there plans for an official Granite 4.2 3B ONNX/WebGPU release?" posted the day after launch (github.com), an indication that browser and edge-runtime support was not fully resolved at release. A developer tutorial published on the community platform Dev.to described Granite 4.2 as representing "a significant leap forward in enterprise-grade, open-source large language models" and concluded it "proves to be an exceptionally competent foundation model for agentic applications" (dev.to) (dev.to); this is a single community author's assessment and should be read as developer sentiment rather than an independent benchmark result.

Market Position and Competitive Landscape

Granite 4.2 competes in a crowded field of open-weight, enterprise-relevant language models, most of which have also converged on permissive or semi-open licensing since 2025. **Meta's Llama 4**, released April 5, 2025 under the custom Llama 4 Community License Agreement, uses a "mixture-of-experts (MoE) architecture" with early fusion for native multimodality; its Maverick variant has 17 billion active parameters, 128 experts, and 400 billion total parameters ([Meta](#)). **Alibaba's Qwen3**, announced April 29, 2025 under the tagline "Qwen3: Think Deeper, Act Faster" (qwenlm.github.io), released its flagship MoE variant with "235 billion total parameters and 22 billion activated" under Apache 2.0, alongside smaller dense models more directly comparable in scale to Granite 4.2. **Mistral AI's Mistral Large 3**, announced December 2, 2025 as part of the "Mistral 3" family (mistral.ai), is "a sparse mixture-of-experts [model] trained with 41B active and 675B total parameters," released under Apache 2.0 (mistral.ai) and supporting a 256,000-token context window per its Hugging Face card (huggingface.co).

Table 2 below summarizes the primary-source specifications of Granite 4.2 alongside these three contemporaries. Note that the table intentionally omits head-to-head benchmark scores across vendors, because no single source reviewed in this research published a directly comparable, apples-to-apples benchmark run across all four families on identical test conditions; publishing invented cross-vendor scores would violate this report's sourcing standard.

Model family	Architecture	Largest / smallest published sizes	Context window	License	Release date
IBM Granite 4.2	Dense decoder-only reasoning transformer	3B to 30B	128K native / 512K extended (3B, 8B, 30B)	Apache 2.0	August 25, 2026

Model family	Architecture	Largest / smallest published sizes	Context window	License	Release date
Meta Llama 4	Mixture-of-experts, native multimodal	~17B active / up to ~400B total	Up to 1M+ tokens (variant-dependent)	Custom Llama 4 Community License	April 5, 2025
Alibaba Qwen3	Mixture-of-experts and dense variants	Dense 0.6B to 32B	Up to 128K (larger dense/MoE variants)	Apache 2.0	April 29, 2025
Mistral Large 3	Sparse mixture-of-experts	41B active / 675B total	256K	Apache 2.0	December 2, 2025

Read together, the table shows that Qwen3's published dense lineup spans 0.6B through 32B. The published parameter ranges therefore do not support ranking Granite 4.2 as the smallest-footprint family overall. Granite 4.2 may still be relevant for organizations prioritizing on-premises or edge portability, IBM enterprise support, and architectural portability.

Data Analysis and Evidence

IBM's Granite team published a benchmark table for Granite 4.2 covering agentic coding, general agentic and tool use, reasoning, chat and instruction following, and long context across its three sizes (huggingface.co). On the **Berkeley Function-Calling Leaderboard (BFCL v4)**, a tool-calling benchmark, IBM reports scores of 52.41 (3B), 50.29 (8B), and 61.39 (30B), notable in that the 8B model scores slightly below the 3B on this specific test. On **AIME25**, a competition-math reasoning benchmark, scores rise consistently with size: 78.33, 86.67, and 89.17. On **GPQA**, a graduate-level science question set, scores are 54.80, 64.14, and 66.41, and on **MMLU-Pro**, a broad knowledge benchmark, scores are 67.84, 74.04, and 77.60 (figures per the same source above).

On coding tasks, **LiveCodeBench v6** scores are 69.71, 73.24, and 75.77 across the 3B/8B/30B sizes (per the same IBM benchmark table). The agentic-coding benchmark **SWE-Bench Verified** is not evaluated for the 3B model at all (reported as "NA"), reflecting that model's exclusion from the agentic-RL training stage, while the 8B and 30B score 47.67 and 57.00 respectively; the same pattern holds for **Terminal-Bench 2.1**, where the 8B and 30B score 20.56 and 29.24. On long-context retrieval, **RULER at 128K tokens** scores 55.30, 71.41, and 81.38, while **RULER at 64K** scores 67.52, 80.99, and 89.96, with retrieval accuracy declining at the longer window across all three sizes, an expected pattern for long-context tasks generally. On instruction-following, **IFBench (prompt-level)** scores are 74.33, 79.33, and 77.17, the only reasoning-family metric in IBM's table where the 30B model scores below the 8B, and on **Arena-Hard-V2**, a chat-quality benchmark, scores are 34.96, 65.19, and 67.93 (all figures in this paragraph per the same IBM technical blog cited above).

Table 3 consolidates these vendor-reported figures.

Benchmark	3B	8B	30B	What it measures
BFCL v4 (tool calling; IBM-reported)	52.41	50.29	61.39	Function/tool-call accuracy
AIME25 (math)	78.33	86.67	89.17	Competition mathematics, pass@1
GPQA (science Q&A)	54.80	64.14	66.41	Graduate-level science reasoning

Benchmark	3B	8B	30B	What it measures
MMLU-Pro (knowledge)	67.84	74.04	77.60	Broad multi-domain knowledge
LiveCodeBench v6 (coding)	69.71	73.24	75.77	Code generation
SWE-Bench Verified (agentic coding)	NA	47.67	57.00	Real repository bug-fixing
Terminal-Bench 2.1 (agentic)	NA	20.56	29.24	Terminal/CLI task completion
RULER 128K (long context)	55.30	71.41	81.38	Long-context retrieval accuracy
IFBench prompt-level	74.33	79.33	77.17	Instruction-following precision
Arena-Hard-V2 (chat)	34.96	65.19	67.93	Open-ended chat quality

All figures in Table 3 are vendor-reported by IBM on its own Hugging Face technical blog, not independently reproduced by a third party for this report, and should be read as claimed rather than externally audited performance. IBM's technical post states that its full benchmark table is presented below. A separate aggregator, **BenchLM**, independently tracked the hosted Granite 4.2 8B route at 144 tokens per second output speed against a field median of 92 tokens per second ([benchlm.ai](#)), and assigned the same model an aggregate "Capability" index of 42 out of 100 against a field median of 58.1 ([benchlm.ai](#)); this composite index is BenchLM's own weighting of vendor-reported per-benchmark scores rather than a new evaluation run, and readers should treat it as one aggregator's relative ranking, not an independent capability measurement in its own right.

Implications and Future Directions

Granite 4.2's release illustrates a broader pattern across the open-weight model market in 2025 and 2026: vendors have converged, within roughly a year, on Apache 2.0 as the default license for enterprise-oriented model families, as seen across IBM Granite 4.2, Alibaba's Qwen3, and Mistral AI's Mistral 3 family alike ([github.com](#)) ([qwenlm.github.io](#)) ([mistral.ai](#)), leaving Meta's custom Llama Community License as a comparative outlier among the four families surveyed. This shift means the practical differentiator between competing open-weight models is increasingly architecture, deployment footprint, and vendor support terms, rather than licensing cost, which is consistent with IBM's own framing of Granite 4.2's dense architecture as a deliberate choice for broad deployment compatibility (per the Version History section above).

Enterprise teams should still cross-reference official model documentation when making deployment decisions, particularly when requirements depend on a specific serving stack or checkpoint.

Conclusion

IBM Granite 4.2 is a dense, decoder-only, Apache 2.0-licensed reasoning model family released August 25, 2026 in 3B, 8B, and 30B parameter sizes, distinguished from its Granite 4.0 predecessor by abandoning that generation's hybrid Mamba-2/transformer Mixture-of-Experts design in favor of the dense architecture IBM introduced with Granite 4.1 and then extended with native chain-of-thought reasoning and reasoning-augmented tool calling. On IBM's own reported benchmarks, capability scales consistently with parameter count on most reasoning and coding

tasks, though not uniformly: the 8B model slightly underperforms the 3B on the BFCL v4 tool-calling benchmark and slightly outperforms the 30B on IFBench instruction-following, evidence that size alone does not predict performance on every task type.

Deployment is well documented across vLLM, SGLang, Ollama, LM Studio, Docker, and IBM's own watsonx.ai platform, though IBM has not published explicit minimum hardware requirements for any size. Licensing is unambiguous and permissive: Apache 2.0 across all three sizes. IBM documentation states: "Only IBM foundation models that you access from watsonx.ai are indemnified by IBM" ([IBM documentation](#)). Pricing is less clear: as of the September 5, 2026 access date, no Granite-4.2-specific watsonx.ai per-token rate was found published, a gap that organizations should close directly with IBM before finalizing hosted-deployment cost estimates rather than relying on prior-generation pricing as a proxy.

In this comparison, [Qwen3's published dense lineup includes 0.6B and 1.7B models](#). For enterprises, particularly in regulated sectors, evaluating Granite 4.2 for agentic, coding, or document-processing workloads, this report's evidence supports treating the model as a credible, well-documented, and permissively licensed option, while recommending independent verification of vendor-reported benchmarks and direct confirmation of current watsonx.ai pricing before committing to a production deployment.

Frequently Asked Questions (FAQs)

What is IBM Granite 4.2?

It is IBM's family of open-weight, dense, decoder-only reasoning language models released August 25, 2026 in 3B, 8B, and 30B parameter sizes under the Apache 2.0 license ([github.com](#)).

How is Granite 4.2 different from Granite 4.0?

Granite 4.0 used a hybrid Mamba-2/transformer Mixture-of-Experts architecture, per IBM's October 2025 announcement cited in the Version History section above; Granite 4.2 is fully dense and decoder-only, and adds explicit chain-of-thought reasoning that neither 4.0 nor the intermediate 4.1 release included, per IBM's own technical blog cited in the Version History section above.

What are the hardware requirements to deploy Granite 4.2?

IBM has not published explicit minimum GPU or VRAM figures for any of the three sizes; documented deployment paths include vLLM (v0.20+), SGLang (v0.5.18+), Ollama, LM Studio, and Docker, with quantized FP8, NVFP4, MXFP4, and GGUF variants available for lower-memory deployment. Context planning should distinguish the 128K native window available across the family from the documented 512K long-context extensions for the [3B](#), [8B](#), and [30B](#) checkpoints.

How is IBM Granite 4.2 deployed?

Options documented by IBM include self-hosting via vLLM or SGLang for production-grade serving, quick local testing via Ollama, LM Studio, or Docker Model Runner, and hosted access via IBM's watsonx.ai platform or partner clouds such as AWS ([ibm.com](#)).

What is the pricing and licensing for Granite 4.2?

The model weights are available under Apache 2.0; commercial use and redistribution remain subject to that license's terms and conditions ([Apache License 2.0](#)). IBM's public [watsonx.ai](#) pricing page says that IBM foundation models include pay-as-you-go pricing per million tokens and hourly rates for on-demand model hosting and deployment ([IBM watsonx.ai pricing](#)).

What are Granite 4.2's enterprise use cases?

IBM positions it for agentic software engineering, terminal/DevOps automation, tool calling, long-document retrieval-augmented generation, and multilingual dialog across financial services, healthcare, telecom, and public-sector contexts ([www.marktechpost.com](#)).

How does Granite 4.2 compare to Granite 4.0 on benchmarks?

No IBM source publishes a direct, side-by-side benchmark table comparing Granite 4.2 against Granite 4.0 on identical tests; IBM's public statements about generational improvement are qualitative rather than a reproducible numeric comparison, a gap this report notes rather than fills with an estimate.

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