

AI's New Frontier: Real-Time Voice, 6× Efficiency, and Hybrid Strategies

Executive Summary

In the past week, the AI frontier has advanced on multiple fronts. Google rolled out Gemini 3.1 Flash Live—a voice-first model enabling faster, natural multilingual conversations (with extended context and live image-query support) ([www.androidcentral.com \[1\]](https://www.androidcentral.com/apps-software/ai/gemini-3-1-flash-live-is-a-massive-boon-to-the-ai-real-time-assistance-for-you-and-me#:~:text=questions%20and%20more%20complex%20topics,conversation%20for%20twice%20as%20long)) ([www.androidcentral.com \[2\]](https://www.androidcentral.com/apps-software/ai/gemini-3-1-flash-live-is-a-massive-boon-to-the-ai-real-time-assistance-for-you-and-me#:~:text=As%20Google%20boosts%20the%20voice,to%20give%20you%20the%20details))—and Google's research team introduced TurboQuant, a compression technique that reduces LLM memory by ~6× (boosting inference ~8×) with no accuracy loss ([www.tomshardware.com \[3\]](https://www.tomshardware.com/tech-industry/artificial-intelligence/googles-turboquant-compresses-llm-kv-caches-to-3-bits-with-no-accuracy-loss#:~:text=Google%27s%20TurboQuant%20reduces%20AI%20LLM,Nvidia%20H100%20GPUs%2C%20compresses%20KV)). Industry leaders are now focusing on hybrid strategies—using both open-source and proprietary models in orchestration systems ([www.pcgamer.com \[4\]](https://www.pcgamer.com/software/ai/nvidia-ceo-jensen-huang-says-open-is-not-a-thing-its-proprietary-and-open-and-sparks-talk-of-ai-orchestras/#:~:text=Cursor%20CEO%20Michael%20Truell%20explains,Jensen%20agrees)) ([www.pcgamer.com \[5\]](https://www.pcgamer.com/software/ai/nvidia-ceo-jensen-huang-says-open-is-not-a-thing-its-proprietary-and-open-and-sparks-talk-of-ai-orchestras/#:~:text=Jensen%20agrees%3A%20,to%20a%20single%20AI%20model))—and many businesses are replacing pricey SaaS tools with custom AI agent systems ([www.techradar.com \[6\]](https://www.techradar.com/pro/how-i-built-an-ai-operating-system-to-run-my-publishing-company#:~:text=A%20%24350%2C000%20Salesforce%20contract%20was,to%20build%20more%20this%20year)) ([www.techradar.com \[7\]](https://www.techradar.com/pro/how-i-built-an-ai-operating-system-to-run-my-publishing-company#:~:text=For%20us%20at%20Man%20of,The%20question%20is%20not)).

References:

- [1] [www.androidcentral.com](https://www.androidcentral.com/apps-software/ai/gemini-3-1-flash-live-is-a-massive-boon-to-the-ai-real-time-assistance-for-you-and-me#:~:text=questions%20and%20more%20complex%20topics,conversation%20for%20twice%20as%20long) — <https://www.androidcentral.com/apps-software/ai/gemini-3-1-flash-live-is-a-massive-boon-to-the-ai-real-time-assistance-for-you-and-me#:~:text=questions%20and%20more%20complex%20topics,conversation%20for%20twice%20as%20long>
- [2] [www.androidcentral.com](https://www.androidcentral.com/apps-software/ai/gemini-3-1-flash-live-is-a-massive-boon-to-the-ai-real-time-assistance-for-you-and-me#:~:text=As%20Google%20boosts%20the%20voice,to%20give%20you%20the%20details) — <https://www.androidcentral.com/apps-software/ai/gemini-3-1-flash-live-is-a-massive-boon-to-the-ai-real-time-assistance-for-you-and-me#:~:text=As%20Google%20boosts%20the%20voice,to%20give%20you%20the%20details>
- [3] [www.tomshardware.com](https://www.tomshardware.com/tech-industry/artificial-intelligence/googles-turboquant-compresses-llm-kv-caches-to-3-bits-with-no-accuracy-loss#:~:text=Google%27s%20TurboQuant%20reduces%20AI%20LLM,Nvidia%20H100%20GPUs%2C%20compresses%20KV) — <https://www.tomshardware.com/tech-industry/artificial-intelligence/googles-turboquant-compresses-llm-kv-caches-to-3-bits-with-no-accuracy-loss#:~:text=Google%27s%20TurboQuant%20reduces%20AI%20LLM,Nvidia%20H100%20GPUs%2C%20compresses%20KV>
- [4] [www.pcgamer.com](https://www.pcgamer.com/software/ai/nvidia-ceo-jensen-huang-says-open-is-not-a-thing-its-proprietary-and-open-and-sparks-talk-of-ai-orchestras/#:~:text=Cursor%20CEO%20Michael%20Truell%20explains,Jensen%20agrees) — <https://www.pcgamer.com/software/ai/nvidia-ceo-jensen-huang-says-open-is-not-a-thing-its-proprietary-and-open-and-sparks-talk-of-ai-orchestras/#:~:text=Cursor%20CEO%20Michael%20Truell%20explains,Jensen%20agrees>
- [5] [www.pcgamer.com](https://www.pcgamer.com/software/ai/nvidia-ceo-jensen-huang-says-open-is-not-a-thing-its-proprietary-and-open-and-sparks-talk-of-ai-orchestras/#:~:text=Jensen%20agrees%3A%20,to%20a%20single%20AI%20model) — <https://www.pcgamer.com/software/ai/nvidia-ceo-jensen-huang-says-open-is-not-a-thing-its-proprietary-and-open-and-sparks-talk-of-ai-orchestras/#:~:text=Jensen%20agrees%3A%20,to%20a%20single%20AI%20model>
- [6] [www.techradar.com](https://www.techradar.com/pro/how-i-built-an-ai-operating-system-to-run-my-publishing-company#:~:text=A%20%24350%2C000%20Salesforce%20contract%20was,to%20build%20more%20this%20year) — <https://www.techradar.com/pro/how-i-built-an-ai-operating-system-to-run-my-publishing-company#:~:text=A%20%24350%2C000%20Salesforce%20contract%20was,to%20build%20more%20this%20year>
- [7] [www.techradar.com](https://www.techradar.com/pro/how-i-built-an-ai-operating-system-to-run-my-publishing-company#:~:text=For%20us%20at%20Man%20of,The%20question%20is%20not) — <https://www.techradar.com/pro/how-i-built-an-ai-operating-system-to-run-my-publishing-company#:~:text=For%20us%20at%20Man%20of,The%20question%20is%20not>

Voice and Multimodal AI Advances

Google has significantly upgraded its foundation models for interactive use. The new Gemini 3.1 Flash Live variant is a voice-optimized model designed for real-time conversations. Google reports it answers queries broadly faster and can maintain the thread of a conversation “twice as long” as previous versions ([www.androidcentral.com \[1\]](https://www.androidcentral.com/apps-software/ai/gemini-3-1-flash-live-is-a-massive-boon-to-the-ai-real-time-assistance-for-you-and-me#:~:text=questions%20and%20more%20complex%20topics,conversation%20for%20twice%20as%20long)). It has been tuned for speech, improving tonal and acoustic nuance detection so that responses sound more natural even in noisy environments ([www.androidcentral.com \[2\]](https://www.androidcentral.com/apps-software/ai/gemini-3-1-flash-live-is-a-massive-boon-to-the-ai-real-time-assistance-for-you-and-me#:~:text=As%20Google%20boosts%20the%20voice,to%20give%20you%20the%20details)). Benchmark early reports suggest Gemini 3.1 scores very high on performance tests, indicating enterprise voice assistants and chatbots will handle speech input much more reliably ([www.androidcentral.com \[3\]](https://www.androidcentral.com/apps-software/ai/gemini-3-1-flash-live-is-a-massive-boon-to-the-ai-real-time-assistance-for-you-and-me#:~:text=As%20Google%20boosts%20the%20voice,to%20give%20you%20the%20details)).

Simultaneously, Gemini is becoming truly multimodal. In the 3.1 update, users can share their camera or screen with the AI, allowing Gemini to answer questions about live images ([www.androidcentral.com \[4\]](https://www.androidcentral.com/apps-software/ai/gemini-3-1-flash-live-is-a-massive-boon-to-the-ai-real-time-assistance-for-you-and-me#:~:text=As%20Google%20boosts%20the%20voice,to%20give%20you%20the%20details)). For example, a field worker could show a product or piece of equipment to their phone camera and ask the AI for details on-the-fly. Google has also incorporated

live translation features (e.g. voice-to-voice translation while streaming media). Overall, Google is pushing Gemini into “voice-first” and vision-assisted use cases.

For executives, these advances mean AI interfaces are moving beyond static text. Expect products with voice-driven workflows (virtual assistants, call centers) and AR/VR apps that listen and look at the world. Customers might soon converse naturally with AI and even get immediate visual assistance. Organizations should prepare to integrate these capabilities: for example, adding voice interfaces to customer support or deploying smart cameras that feed Gemini data. The path is clear: real-time, spoken and visual interaction is the new baseline for AI competencies, and companies need to plan for products and workflows that leverage voice+vision AI.

References:

- [1] [www.androidcentral.com — https://www.androidcentral.com/apps-software/ai/gemini-3-1-flash-live-is-a-massive-boon-to-the-ai-real-time-assistance-for-you-and-me#:~:text=questions%20and%20more%20complex%20topics,conversation%20for%20twice%20as%20long](https://www.androidcentral.com/apps-software/ai/gemini-3-1-flash-live-is-a-massive-boon-to-the-ai-real-time-assistance-for-you-and-me#:~:text=questions%20and%20more%20complex%20topics,conversation%20for%20twice%20as%20long)
- [2] [www.androidcentral.com — https://www.androidcentral.com/apps-software/ai/gemini-3-1-flash-live-is-a-massive-boon-to-the-ai-real-time-assistance-for-you-and-me#:~:text=real,as%20the%20ability%20to%20recognize](https://www.androidcentral.com/apps-software/ai/gemini-3-1-flash-live-is-a-massive-boon-to-the-ai-real-time-assistance-for-you-and-me#:~:text=real,as%20the%20ability%20to%20recognize)
- [3] [www.androidcentral.com — https://www.androidcentral.com/apps-software/ai/gemini-3-1-flash-live-is-a-massive-boon-to-the-ai-real-time-assistance-for-you-and-me#:~:text=real,as%20the%20ability%20to%20recognize](https://www.androidcentral.com/apps-software/ai/gemini-3-1-flash-live-is-a-massive-boon-to-the-ai-real-time-assistance-for-you-and-me#:~:text=real,as%20the%20ability%20to%20recognize)
- [4] [www.androidcentral.com — https://www.androidcentral.com/apps-software/ai/gemini-3-1-flash-live-is-a-massive-boon-to-the-ai-real-time-assistance-for-you-and-me#:~:text=As%20Google%20boosts%20the%20voice,to%20give%20you%20the%20details](https://www.androidcentral.com/apps-software/ai/gemini-3-1-flash-live-is-a-massive-boon-to-the-ai-real-time-assistance-for-you-and-me#:~:text=As%20Google%20boosts%20the%20voice,to%20give%20you%20the%20details)

Efficiency Gains: TurboQuant & Compute Costs

Google has also announced a breakthrough in model efficiency called TurboQuant. This training-free algorithm dramatically compresses the key-value memory caches that LLMs use for attention. In benchmark tests (e.g. on Nvidia H100 GPUs), TurboQuant shrank these caches by roughly 6× and accelerated attention computations up to 8×, all with zero loss in model accuracy (www.tomshardware.com [1]). In practical terms, LLM inference that once needed huge GPU memory can now run on much smaller hardware with the same performance.

The business impact is significant. Memory and compute cost have been major barriers to deploying large-context LLMs at scale. TurboQuant’s “6× smaller, 8× faster” result suggests both cloud bills and on-prem hardware costs can drop sharply. Enterprises could run bigger models (longer chats, bigger datasets) on existing servers or reduce billable inference time. This could lower the price of AI services: for example, a SaaS LLM provider might cite lower runtime costs when negotiating contracts.

Going forward, executives should watch how quickly these techniques spread. Google says TurboQuant will be open-sourced, so model platforms and libraries (e.g. Hugging Face, LangChain, Nvidia stacks) may incorporate such compression soon. If widely adopted, the effective dollar-per-token for AI tasks will fall. Strategically, this means companies can plan for richer AI applications (more users, longer inputs) without linear cost increases. In short, expect the efficiency frontier to move markedly forward into the coming year, enabling higher-volume AI services for the same budget.

References:

- [1] [www.tomshardware.com — https://www.tomshardware.com/tech-industry/artificial-intelligence/googles-turboquant-compresses-llm-kv-caches-to-3-bits-with-no-accuracy-loss#:~:text=Google%27s%20TurboQuant%20reduces%20AI%20LLM,Nvidia%20H100%20GPUs%2C%20compresses%20KV](https://www.tomshardware.com/tech-industry/artificial-intelligence/googles-turboquant-compresses-llm-kv-caches-to-3-bits-with-no-accuracy-loss#:~:text=Google%27s%20TurboQuant%20reduces%20AI%20LLM,Nvidia%20H100%20GPUs%2C%20compresses%20KV)

The Open/Closed Divide and AI Orchestration

Industry thinking is also shifting on model strategy. At Nvidia’s recent conference, CEO Jensen Huang argued that the old “open vs proprietary” model distinction is already outdated (www.pcgamer.com [1]). The emerging concept is AI orchestration: systems that coordinate multiple models to solve complex tasks. For instance, one panel described future agents as compound systems where you delegate tasks and an AI “conductor” assigns them to specialist “sub-agent” models like musicians playing different instruments (www.pcgamer.com [2]). In this view, a user’s request might be handled

by an open-source model for general knowledge and by a proprietary model for premium or sensitive work.

This hybrid mindset reflects confidence in open models. Reflection AI's CEO Misha Laskin points out that there's "nothing fundamentally different between an open and a closed model" ([www.pcgamer.com \[3\]](https://www.pcgamer.com/3/)), meaning top-tier capabilities exist on both sides. Jensen Huang concurs that even in a company with a private LLM, open-source models will still be used for broad tasks, with the in-house model reserved as the "crown jewel" specialist ([www.pcgamer.com \[4\]](https://www.pcgamer.com/4/)). The implication: businesses should not view open-source as inherently inferior. Instead, they should expect to mix open and closed models depending on cost, performance, and data needs.

Practically speaking, this means building flexible AI pipelines. Organizations will likely use an open model via API or on-prem for general tasks (like summarization or common queries) and keep proprietary models for domain-specific work. Toolchains (like LangChain, Ray Serve, or upcoming open orchestrators) are emerging to support these multi-model workflows. Strategy-wise, executives should treat AI as a portfolio of assets. As one analogy suggests, consider the multiple AIs in your system as instruments and your products as compositions they create. Planning should involve integration layers that route tasks to the right model, rather than a single vendor lock-in.

References:

- [1] [www.pcgamer.com — https://www.pcgamer.com/software/ai/nvidia-ceo-jensen-huang-says-open-is-not-a-thing-its-proprietary-and-open-and-sparks-talk-of-ai-orchestras/#:~:text=Nvidia%20CEO%20Jensen%20Huang%20says,Mar%2026%2016%3A25%3A14%202026%20UTC](https://www.pcgamer.com/software/ai/nvidia-ceo-jensen-huang-says-open-is-not-a-thing-its-proprietary-and-open-and-sparks-talk-of-ai-orchestras/#:~:text=Nvidia%20CEO%20Jensen%20Huang%20says,Mar%2026%2016%3A25%3A14%202026%20UTC)
- [2] [www.pcgamer.com — https://www.pcgamer.com/software/ai/nvidia-ceo-jensen-huang-says-open-is-not-a-thing-its-proprietary-and-open-and-sparks-talk-of-ai-orchestras/#:~:text=Cursor%20CEO%20Michael%20Truell%20explains,Jensen%20agrees](https://www.pcgamer.com/software/ai/nvidia-ceo-jensen-huang-says-open-is-not-a-thing-its-proprietary-and-open-and-sparks-talk-of-ai-orchestras/#:~:text=Cursor%20CEO%20Michael%20Truell%20explains,Jensen%20agrees)
- [3] [www.pcgamer.com — https://www.pcgamer.com/software/ai/nvidia-ceo-jensen-huang-says-open-is-not-a-thing-its-proprietary-and-open-and-sparks-talk-of-ai-orchestras/#:~:text=all%20these%20models%20be%20open,open%20and%20a%20closed%20model](https://www.pcgamer.com/software/ai/nvidia-ceo-jensen-huang-says-open-is-not-a-thing-its-proprietary-and-open-and-sparks-talk-of-ai-orchestras/#:~:text=all%20these%20models%20be%20open,open%20and%20a%20closed%20model)
- [4] [www.pcgamer.com — https://www.pcgamer.com/software/ai/nvidia-ceo-jensen-huang-says-open-is-not-a-thing-its-proprietary-and-open-and-sparks-talk-of-ai-orchestras/#:~:text=Jensen%20agrees%3A%20,to%20a%20single%20AI%20model](https://www.pcgamer.com/software/ai/nvidia-ceo-jensen-huang-says-open-is-not-a-thing-its-proprietary-and-open-and-sparks-talk-of-ai-orchestras/#:~:text=Jensen%20agrees%3A%20,to%20a%20single%20AI%20model)

Enterprise Strategy: Custom Agents vs SaaS

The business world is already restructuring around AI. Recent reports call it a "SaaS-pocalypse": legacy SaaS firms have lost huge market value as AI automates their key functions. In practice, companies find that one AI agent can do the work of dozens of specialized apps. For example, after building an internal AI system, one publisher canceled a \$350K Salesforce contract ([www.techradar.com \[1\]](https://www.techradar.com/1/)). Industry survey data backs this up: 35% of enterprises have already replaced at least one SaaS tool with custom AI, and 78% plan more such builds this year ([www.techradar.com \[2\]](https://www.techradar.com/2/)). In other words, many companies are opting to build their own AI-driven solutions instead of buying new subscriptions.

A concrete case illustrates the value. Man of Many, an Australian media firm, engineered an AI "operating system" named Otto. It was built using Anthropic's Claude (no human coding required) in about a week ([www.techradar.com \[3\]](https://www.techradar.com/3/)). Otto fetches daily revenue and traffic reports, flags issues, compiles editorial briefs, and even drafts standard documents. Its security layer restricts high-risk actions to human approval. Importantly, Otto cost only a fraction of the company's previous SaaS budget ([www.techradar.com \[4\]](https://www.techradar.com/4/)). This shows that with modern LLM tools, small teams can replace dozens of subscriptions and manual tasks, and focus on higher-level strategy.

What does this mean for leaders? Almost every industry faces this shift: for example, 97% of publishers now say back-end AI automation is important ([www.techradar.com \[5\]](https://www.techradar.com/5/)). Executives should map their current tech stack and identify routine workflows ripe for AI. Adopting an "AI-first" approach – letting AI handle data aggregation, scheduling, reporting – can cut costs and boost productivity. It also means re-evaluating vendor roadmaps: if a needed feature is delayed, consider building an

internal agent. The key takeaway is that the capability frontier is moving from new apps to transforming existing processes: companies will either build custom AI agents for tasks or risk being outflanked by competitors who do.

References:

- [1] [www.techradar.com — https://www.techradar.com/pro/how-i-built-an-ai-operating-system-to-run-my-publishing-company#:~:text=The%20reason%20is%20simple,to%20build%20more%20this%20year](https://www.techradar.com/pro/how-i-built-an-ai-operating-system-to-run-my-publishing-company#:~:text=The%20reason%20is%20simple,to%20build%20more%20this%20year)
- [2] [www.techradar.com — https://www.techradar.com/pro/how-i-built-an-ai-operating-system-to-run-my-publishing-company#:~:text=A%20%24350%2C000%20Salesforce%20contract%20was,to%20build%20more%20this%20year](https://www.techradar.com/pro/how-i-built-an-ai-operating-system-to-run-my-publishing-company#:~:text=A%20%24350%2C000%20Salesforce%20contract%20was,to%20build%20more%20this%20year)
- [3] [www.techradar.com — https://www.techradar.com/pro/how-i-built-an-ai-operating-system-to-run-my-publishing-company#:~:text=each%20give%20you%2059,a%20week%20of%20active%20development](https://www.techradar.com/pro/how-i-built-an-ai-operating-system-to-run-my-publishing-company#:~:text=each%20give%20you%2059,a%20week%20of%20active%20development)
- [4] [www.techradar.com — https://www.techradar.com/pro/how-i-built-an-ai-operating-system-to-run-my-publishing-company#:~:text=For%20us%20at%20Man%20of,The%20question%20is%20not](https://www.techradar.com/pro/how-i-built-an-ai-operating-system-to-run-my-publishing-company#:~:text=For%20us%20at%20Man%20of,The%20question%20is%20not)
- [5] [www.techradar.com — https://www.techradar.com/pro/how-i-built-an-ai-operating-system-to-run-my-publishing-company#:~:text=We%20are%20not%20the%20only,end%20AI%20automation%20%22important](https://www.techradar.com/pro/how-i-built-an-ai-operating-system-to-run-my-publishing-company#:~:text=We%20are%20not%20the%20only,end%20AI%20automation%20%22important)

Key Statistics

- Google's TurboQuant compresses LLM caches ~6× smaller and yields ~8× faster inference (no loss of accuracy) ([[www.tomshardware.com](https://www.tomshardware.com/tech-industry/artificial-intelligence/googles-turboquant-compresses-llm-kv-caches-to-3-bits-with-no-accuracy-loss#:~:text=Google%27s%20TurboQuant%20reduces%20AI%20LLM,Nvidia%20H100%20GPUs%2C%20compresses%20KV)])(<https://www.tomshardware.com/tech-industry/artificial-intelligence/googles-turboquant-compresses-llm-kv-caches-to-3-bits-with-no-accuracy-loss#:~:text=Google%27s%20TurboQuant%20reduces%20AI%20LLM,Nvidia%20H100%20GPUs%2C%20compresses%20KV>)).
- 35% of enterprises replaced a SaaS app with custom AI this year, and 78% plan more builds ([[www.techradar.com](https://www.techradar.com/pro/how-i-built-an-ai-operating-system-to-run-my-publishing-company#:~:text=A%20%24350%2C000%20Salesforce%20contract%20was,to%20build%20more%20this%20year)])(<https://www.techradar.com/pro/how-i-built-an-ai-operating-system-to-run-my-publishing-company#:~:text=A%20%24350%2C000%20Salesforce%20contract%20was,to%20build%20more%20this%20year>)).
- 97% of publishers say back-end AI automation is now “important” ([[www.techradar.com](https://www.techradar.com/pro/how-i-built-an-ai-operating-system-to-run-my-publishing-company#:~:text=We%20are%20not%20the%20only,end%20AI%20automation%20%22important)])(<https://www.techradar.com/pro/how-i-built-an-ai-operating-system-to-run-my-publishing-company#:~:text=We%20are%20not%20the%20only,end%20AI%20automation%20%22important>)).

KEY TAKEAWAY

AI's frontier has surged: Google's new Gemini 3.1 powers real-time multilingual voice+vision interactions, and Google's TurboQuant cuts LLM memory needs by roughly 6×. Experts now foresee hybrid AI orchestration (mixing open-source and proprietary models), and many companies are replacing bloated SaaS with custom AI agents.

Sources

- [Google's TurboQuant reduces AI LLM cache memory capacity requirements by at least six timesa0— up to 8x performance boost on Nvidia H100 GPUs, compresses KV caches to 3 bits with no accuracy loss](https://www.tomshardware.com/tech-industry/artificial-intelligence/googles-turboquant-compresses-llm-kv-caches-to-3-bits-with-no-accuracy-loss)
<https://www.tomshardware.com/tech-industry/artificial-intelligence/googles-turboquant-compresses-llm-kv-caches-to-3-bits-with-no-accuracy-loss>
- [Gemini 3.1 Flash Live is a massive boon to the AI's real-time assistance for you and me](https://www.androidcentral.com/apps-software/ai/gemini-3-1-flash-live-is-a-massive-boon-to-the-ais-real-time-assistance-for-you-and-me)
<https://www.androidcentral.com/apps-software/ai/gemini-3-1-flash-live-is-a-massive-boon-to-the-ais-real-time-assistance-for-you-and-me>
- [Nvidia CEO Jensen Huang says of AI that 'open is not a thing, it's proprietary and open', sparking a conversation about bot orchestras](https://www.pcgamer.com/software/ai/nvidia-ceo-jensen-huang-says-open-is-not-a-thing-its-proprietary-and-open-and-sparks-talk-of-ai-orchestras)
<https://www.pcgamer.com/software/ai/nvidia-ceo-jensen-huang-says-open-is-not-a-thing-its-proprietary-and-open-and-sparks-talk-of-ai-orchestras>
- [How I built an AI operating system to run my publishing company](https://www.techradar.com/pro/how-i-built-an-ai-operating-system-to-run-my-publishing-company)
<https://www.techradar.com/pro/how-i-built-an-ai-operating-system-to-run-my-publishing-company>

