

The emma Platform: Operationalizing AI with Confidence

Built-in Agility, Cost Control, and Governance for Enterprise-Grade AI

Key Capabilities for Enterprise-grade AI Projects:

- Multi-cloud coverage
- On-prem & Hybrid cloud support
- Unified TCO dashboard
- Predictive cost engine
- Real-time cost and usage analytics
- Native integration with ROI modellers
- Hard spend caps for projects
- Automatic workload-locality enforcement
- Sovereign by design

Executive Summary

AI has moved past experimentation. CIOs and CTOs are now expected to turn pilots into production-grade systems that deliver measurable value. Yet across industries, the same roadblocks stand in the way:

- Soaring GPU and infrastructure costs & unclear ROI
- Increasing scrutiny around data governance and compliance
- Fragmented cloud & data Environments

What was once an innovation challenge has become an operational one.

emma helps technology leaders cut through this complexity with a unified cloud management layer purpose-built for modern AI workloads. It brings together compute, storage, and network resources from AWS, Azure, GCP, and sovereign clouds into one control plane – so teams can deploy, monitor, and optimize AI infrastructure without losing visibility or control.

The Core Challenge: Operationalizing AI Without Losing Control

AI experimentation often happens in isolated sandboxes or single clouds (e.g., AWS Sagemaker or Azure ML), but production AI means orchestrating massive compute, data, and network resources, often across hyperscalers, neo clouds, and sovereign clouds.

The problem is, each cloud comes with separate cost models, security policies, and network overheads. The result is:

- **Operational sprawl**
- **Uncontrolled budgets**
- **Data residency issues**
- **Slower AI delivery cycles**

And while these challenges are well known, the reality is that organizations don't have the luxury of months (or years) to figure them out. They're under pressure to operationalize AI now – to move from pilot to production to revenue generation as fast as possible, without compromising performance, cost efficiency, or compliance.

emma: The Unified Cloud Management Layer for AI

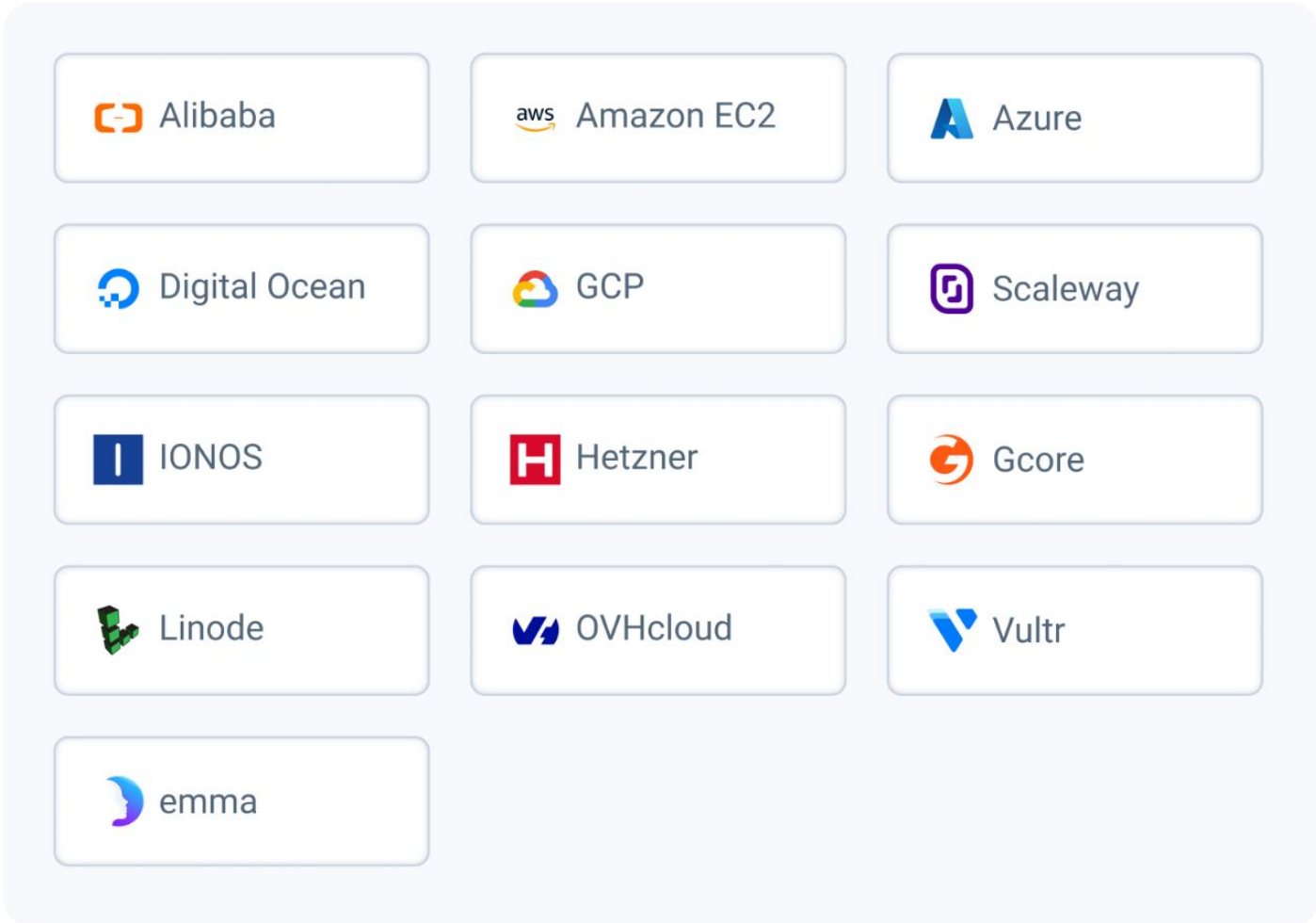
The emma platform expedites the deployment and scaling of enterprise-grade AI by simplifying infrastructure and enabling end-to-end visibility and control, so innovation can move faster and teams can focus on building efficient, sustainable, and compliant AI products that grow ARR.

KEY CAPABILITIES:

1 Infrastructure Freedom & Flexibility

Without true multi-cloud agility, enterprises can find themselves locked out of specialized compute, advanced AI services, and sovereign cloud options from other vendors.

emma lets organizations choose the best-fit environment for every AI workload – public, private (on-prem), or sovereign. It provides a **one-stop platform to choose, manage, and optimize** workloads across all clouds – from **AWS, Azure, and GCP to EU-based sovereign providers** like emma, OVHcloud, and Gcore.

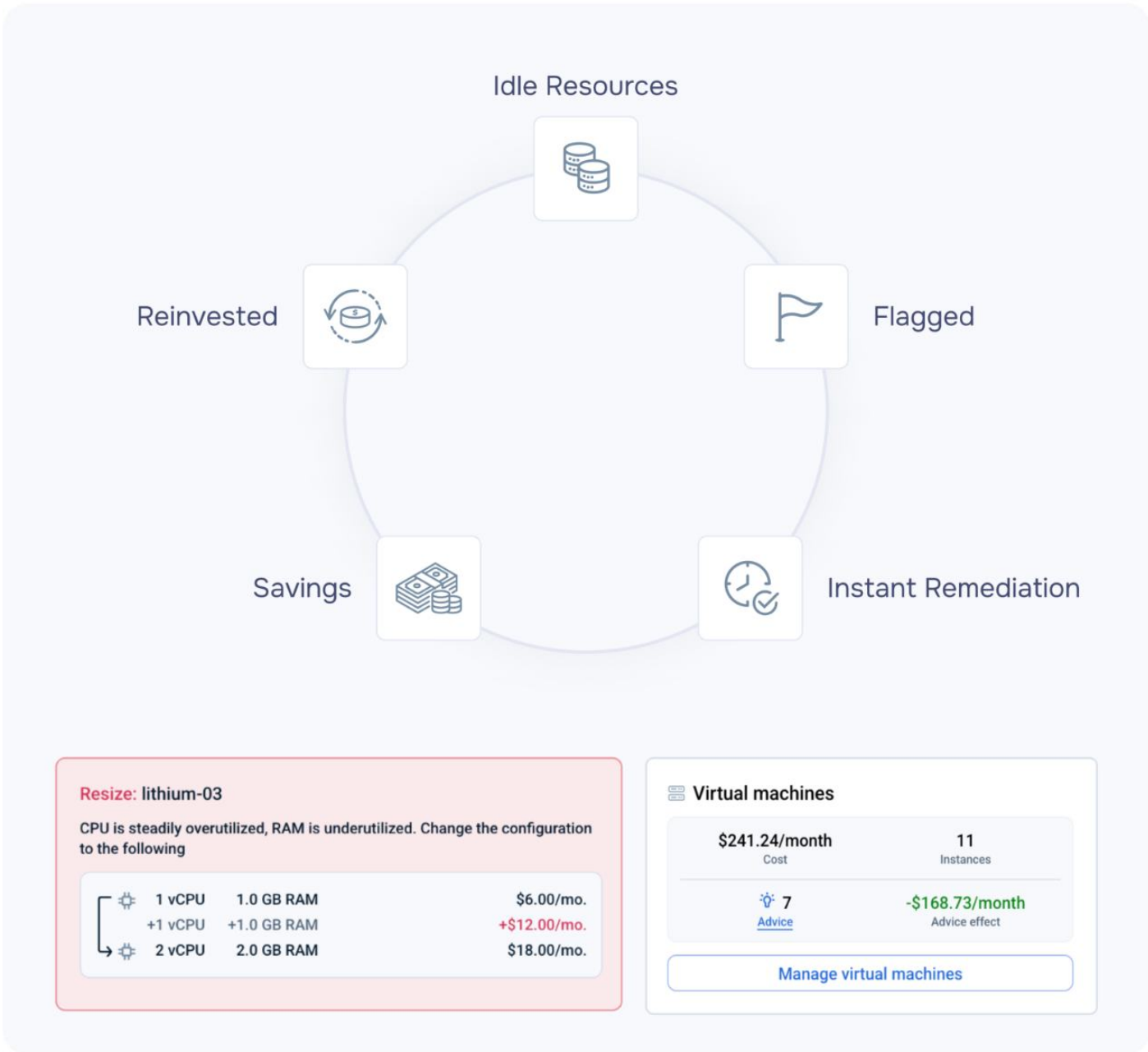


2 Financial Efficiency by Design

Without unified visibility into costs and GPU usage across each cloud and AI project, AI spending can spiral fast. As a result, CIOs and CTOs struggle to balance innovation speed with budget control.

emma offers **real-time cost breakdowns by cloud provider, team, and project**. This helps organizations pinpoint exactly where budgets are being consumed and why. Its **cost expense analysis dashboards** visualize spending trends, real-time anomalies, and allocation ratios, allowing everyone to stay informed and proactive about budgets.

In addition to cost breakdowns and analysis, emma automatically detects and flags **underutilized, over-provisioned, or forgotten** AI resources, such as idle GPU clusters, staging environments left running, or duplicated resources. Built-in **optimization recommendations** let teams **remediate waste in one click**, through rightsizing, scheduling, or decommissioning options.



Beyond operational control, emma integrates with **ROI tools to calculate the real return on AI** investments, helping CFOs and CIOs align budgets to business outcomes. The result is a continuous feedback loop between innovation and accountability, where every GPU hour and every dollar drives measurable value.

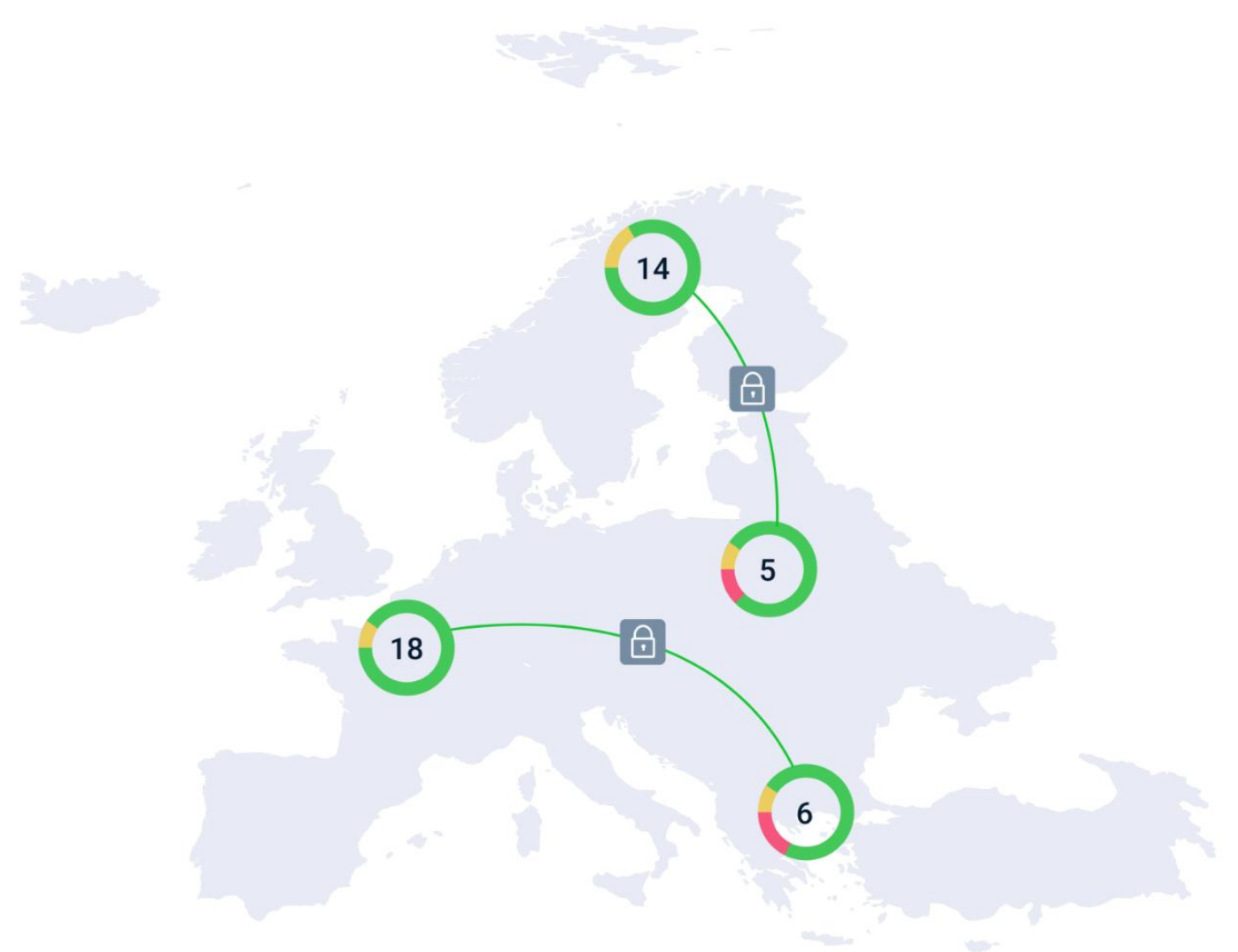
KEY CAPABILITIES:

3 Infrastructure Freedom & Flexibility

Governance isn't an afterthought. emma enforces data residency by region, extends access policies across providers, and maintains auditable trails for both internal and external oversight.

Data is the lifeblood of AI operations, and compliance isn't optional – it's existential. emma builds governance and sovereignty directly into the platform through proactive, user-defined policies and automated enforcement, so organizations can deploy, manage, and scale AI responsibly, regardless of where they operate.

With **region-based data and workload residency enforcement**, emma ensures workloads and data remain within designated geographic and regulatory boundaries, whether hosted on AWS, Azure, GCP, or sovereign providers like emma, OVHcloud, and Gcore. Its policy orchestration engine extends IAM and access controls consistently across all providers, eliminating gaps between environments. This makes it easier to prove and maintain continuous compliance without manual intervention.



And for organizations operating under European or regional sovereignty mandates, emma's **Luxembourg-based foundation, in-house data center with on-demand baremetal and GPU servers, and sovereign-ready networking backbone** provide a regulatory advantage as well. It ensures that AI data and workloads remain both performant and jurisdictionally compliant, both in-transit and at rest.

The Outcome: Scalable, Compliant AI that Delivers Real ROI

With emma, organizations can take ideas from POCs to production faster, at enterprise scale. They can operationalize AI projects confidently by balancing agility, compliance, and cost-efficiency in one platform. No more patchwork of tools, providers, or scripts. No more trade-offs between innovation, governance, and financial control. emma transforms fragmented AI infrastructure into a cohesive, compliant, and cost-optimized foundation built for sustained, enterprise-grade innovation.

- **Achieve true infrastructure freedom.** Run AI workloads where they perform best, whether on hyperscalers, sovereign clouds, or edge nodes. Deploy and optimize workloads across NVIDIA GPUs, custom chips, and regional data centers
- **Gain financial efficiency by design.** Get real-time, unit-level cost breakdowns and AI-powered automations to eliminate waste and enforce spend governance, ensuring predictable ROI and sustainable AI operations.
- **Maintain security and sovereignty while innovating globally.** Leverage workload residency controls and automated enforcement to ensure every AI workload runs in a fully sovereign, compliant environment.

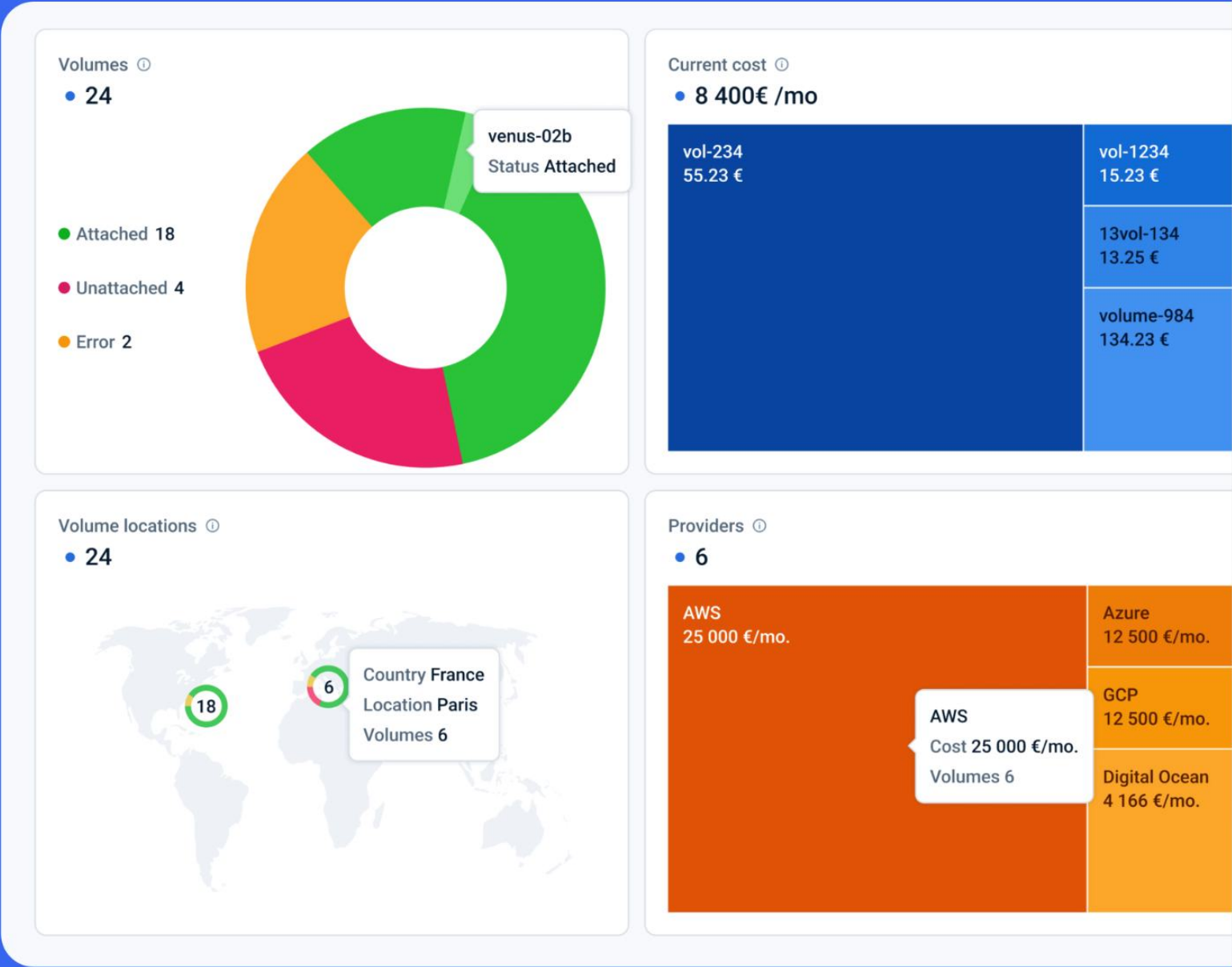
The future of AI won't be defined by who experiments fastest, but by who operationalizes it best. And emma is actively powering that operational excellence.

ABOUT EMMA

emma transforms how the world harnesses the cloud. By removing silos between providers and environments, emma unifies fragmented infrastructures into one seamless experience, unlocking the cloud’s full potential.

With emma’s end-to-end platform, organizations can build, optimize, and govern their cloud environments without lock-in, waste or compromise. Combining intelligent orchestration, proactive optimization, and seamless connectivity into a single platform, emma empowers businesses to move faster, scale smarter, and stay ahead in a rapidly changing world.

Founded in 2021 in Luxembourg, emma has raised \$23M in venture funding from investors including RTP Global, CircleRock Capital, AltaIR Capital, and Smartfin.



Visit emma at:
emma.ms

Contact us:
info@emma.ms

Explore the Product Hub:
emma.ms/product-hub