

The Unglamorous Layer: Event Streaming at 150K Events/Second

Owning the event streaming and persistence platform that feeds ML, ads, and analytics — cutting onboarding from days to minutes and delivering multi-million-dollar annual cloud savings.

Context

Every ML model, personalization system, ad measurement pipeline, and analytics dashboard depends on one thing: clickstream and event data arriving reliably, quickly, and in a usable shape. I owned that platform — the event streaming and persistence layer serving ad tech, customer intelligence, storefront, supply chain, and analytics teams.

The Problem

The platform had grown organically and expensively. Producer onboarding was a days-long, ticket-driven process. Cloud costs were scaling faster than traffic. Reliability issues were hard to debug because failed events simply vanished. None of this was visible to executives until it broke something they cared about — which is precisely what makes data platform work hard to prioritize and easy to underinvest in.

Approach & Key Decisions

- **Self-service as the scaling strategy.** We rebuilt onboarding around a cloud-native pub/sub architecture with authenticated endpoints and infrastructure-as-code (Terraform) provisioning — cutting producer onboarding from days to minutes and removing the platform team from the critical path.
- **Cost as a product metric.** I ran a systematic cloud cost optimization program — cluster right-sizing, retention policy redesign, storage tiering — treating each optimization like a feature with a business case. The program delivered multi-million-dollar annualized savings with further identified runway.
- **Reliability you can debug.** We launched retry and dead-letter-queue capabilities so failed events became inspectable artifacts instead of silent data loss — transforming incident triage and giving downstream ML teams confidence in training data completeness.
- **Observability at the source.** Partnering with ML platform teams, we ensured newly deployed self-managed models shipped with out-of-the-box observability tied back to event data, so model monitoring wasn't bolted on after the first incident.

Impact

Performance	
	Sustained throughput of ~150K events/second with P95 latency under 100 ms.

Cost	Multi-million-dollar annualized cloud savings, with realized gains exceeding the original program target.
Velocity	Producer onboarding reduced from days to minutes via self-serve, IaC-based provisioning.

What I'd Tell Other PMs

“Every AI product stands on a data layer someone has to product-manage. If nobody treats it as a product, you find out during an incident.”

Infrastructure PM work rewards the same instincts as consumer work — remove friction, measure everything, make the right path the easy path. The difference is your wins are invisible when they succeed, so you must learn to narrate value in dollars, latency, and developer-hours.

Sneha Chitte · Senior Product Manager – ML & Personalization · Greater Boston, MA · Metrics are directional and anonymized to protect confidential information.