



서울

From Blackbox To Insight: Observability for AI Agents in Production

Brandon Kang, Principal Technical Solutions Architect
Akamai Technologies

Mostafa Radwan, Senior Solutions Engineer
Datadog



Agenda

- Intro to AI Agents
- What's unique about Agentic Systems?
- AI Agents Observability
- Demo
- Key Takeaways
- Q&A

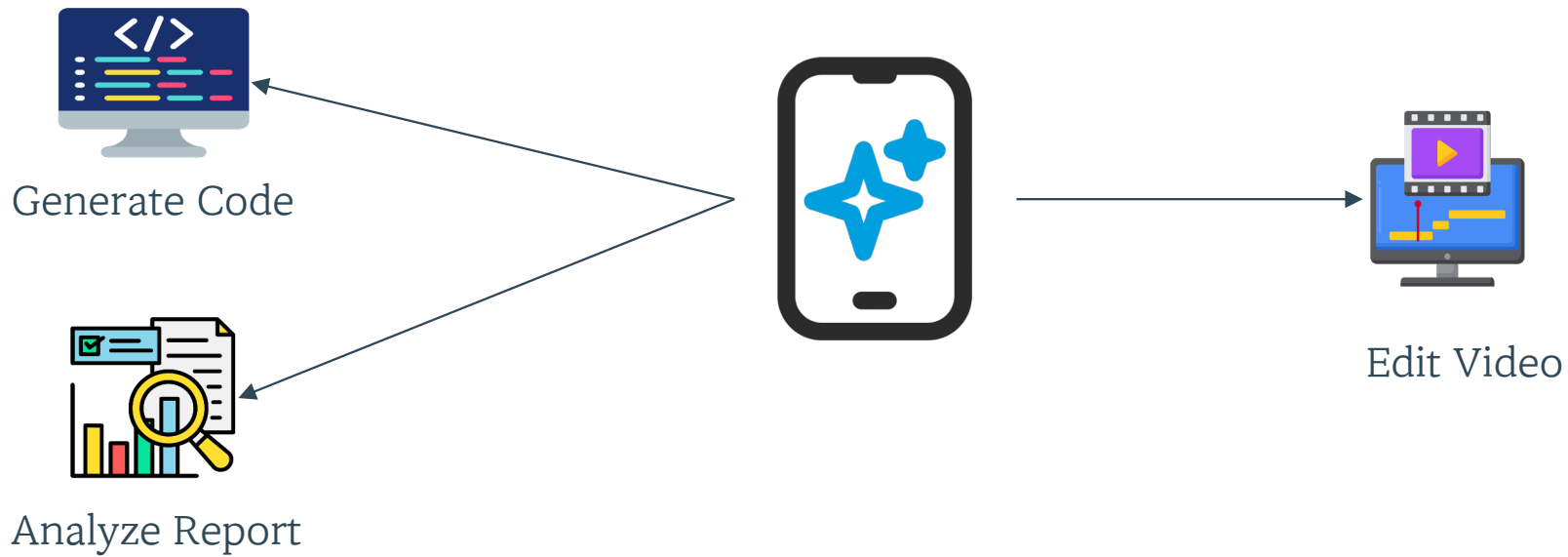
AI Agents Everywhere



“AI agents still need tools to act on the real world”



AI Agents Everywhere



AI Agents Everywhere

Search

FORTUNE

Subscribe

Sign in

Home

Latest

Fortune 500

Finance

Tech

Leadership

Lifestyle

Rankings

Multimedia

New video show that brings you inside the Fortune newsroom and the boardroom

Watch Fortune Daily, our new video show that brings you inside the Fortune newsroom and the b

AI • TECH

Uber burned through its entire 2026 AI budget in four months. Now its COO is questioning whether it's worth it

By Jake Angelo

Former News Fellow

May 26, 2026, 2:03 PM ET

Uber COO Andrew Macdonald

ZED JAMESON/BLOOMBERG VIA GETTY IMAGES



All topics

AXIOS

Search

Together, we reinvented the corporate grind.

Learn More

REINVENTED WITH accenture

15 hours ago - Technology

OpenAI says its AI agents breached its own systems before Hugging Face

Sam Sabin

✉

📞

f

%

in

W

Add Axios on Google



What's unique about Agentic Systems?

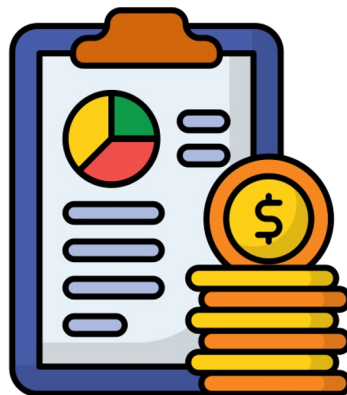
- Non-Deterministic
- Prompt Engineering
- Memory Access (i.e: OpenClaw)
- Tool Calls (aka Integrations)
- LLM tokens and/or GPUs (\$\$\$)



AI Agents Observability



Performance



Cost



Behaviour



AI Agents Observability Tools

Project	Best For	Pros	Cons
OpenTelemetry (OTel)	Vendor-neutral	Instrument once and export to many backends	No UI, backend, or alerting by itself
OpenLLMetry	Instrumentation library and SDK for LLM apps that outputs standard OTel.	Framework/provider instrumentations	Not a full observability backend (limited without dashboards, analysis,..etc)
OpenLIT	OTel-native AI observability that includes cost, evals, prompts, and GPU	OTel-native with auto-instrumentation and various LLM and GPU integrations	An emerging ecosystem



Our Approach: OTel Unified Observability

- Full observability stack for AI agents
- Components:
 -  Prometheus - Metrics collection & storage
 -  Jaeger - Distributed tracing
 -  Grafana - Real-time visualization
 -  Custom Detectors - Anomaly detection
- Example Tech Stack:
 - OpenTelemetry SDK 1.24.0, Prometheus 2.50, Grafana 10.4
 - Python 3.12, Docker Compose
 - NVIDIA RTX Pro 6 000 Blackwell GPU × 4 (96GB VRAM each)



Hallucination Detection: 4-Layer System

- LLMs generate plausible but false information
- Future Date Detection
 - Regex: mentions of 2027+ → 99% hallucination
- Fake Citation Detection
 - Verify against arXiv, CrossRef APIs
 - Example: [Johnson et al., 2027] → fake paper
- Self-Contradiction Detection
 - NLI model (DeBERTa-v3) finds logical conflicts
- Overconfidence Detection
 - Count absolute terms: 'definitely', 'always'
 - Threshold: 3+ → overconfident

```
class HallucinationDetector:
    def _detect_future_dates(self, response):
        pattern = r'\b(202[5-9]|20[3-9]\d)\b'
        matches = re.findall(pattern, response)
        current_year = datetime.now().year

        for year in matches:
            if int(year) > current_year:
                return True # Future date = hallucination!
        return False

    def detect(self, response, context):
        checks = [
            self._detect_future_dates(response),
            self._detect_fake_citations(response),
            self._detect_self_contradiction(response),
            self._detect_overconfidence(response)
        ]
        return any(checks)
```

* Multi-layered detection increases precisions from 60%(single method) to 94%

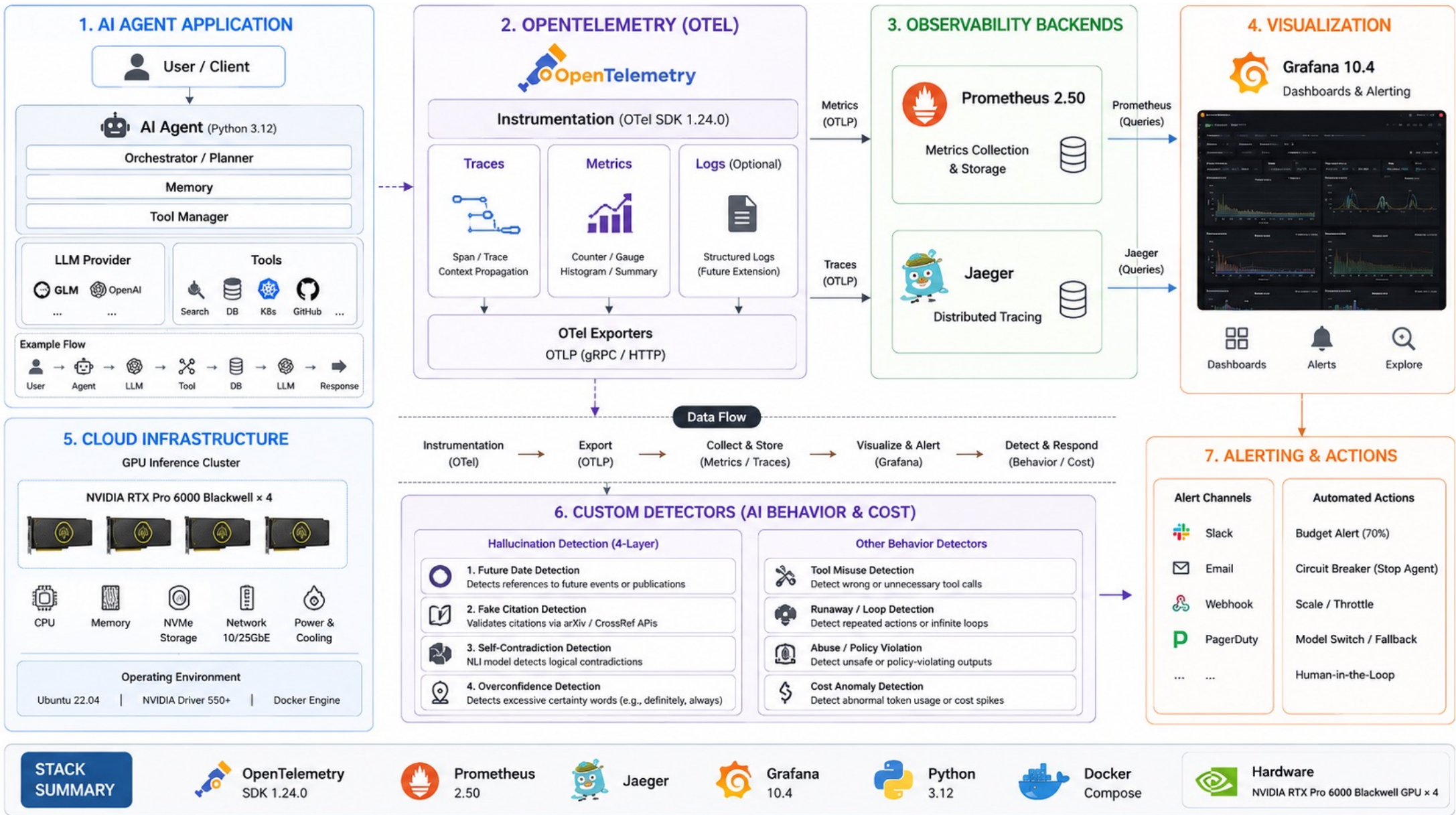
GPU-Based LLM Cost Calculation

- Open-weight models on owned GPUs have real costs!
- Formula:
 - $\text{GPU Hourly Cost} / \text{Tokens Per Hour} = \text{Cost Per Token}$
- Example: GLM-5.2 on NVIDIA RTX Pro 6000 Blackwell
 - • GPU rental: USD 3.5/hour (cloud benchmark)
 - • Throughput: 150 tokens/sec = 540K tokens/Hour
 - • Cost per 1K tokens: \$0.00648
- Our Pricing (with overhead):
 - • GLM-5.2 Input: \$0.008/1K tokens
 - • GLM-5.2 Output: \$0.025/1K (3× slower)
- vs. GPT-4 API: \$0.03/\$0.06 → 82% savings!

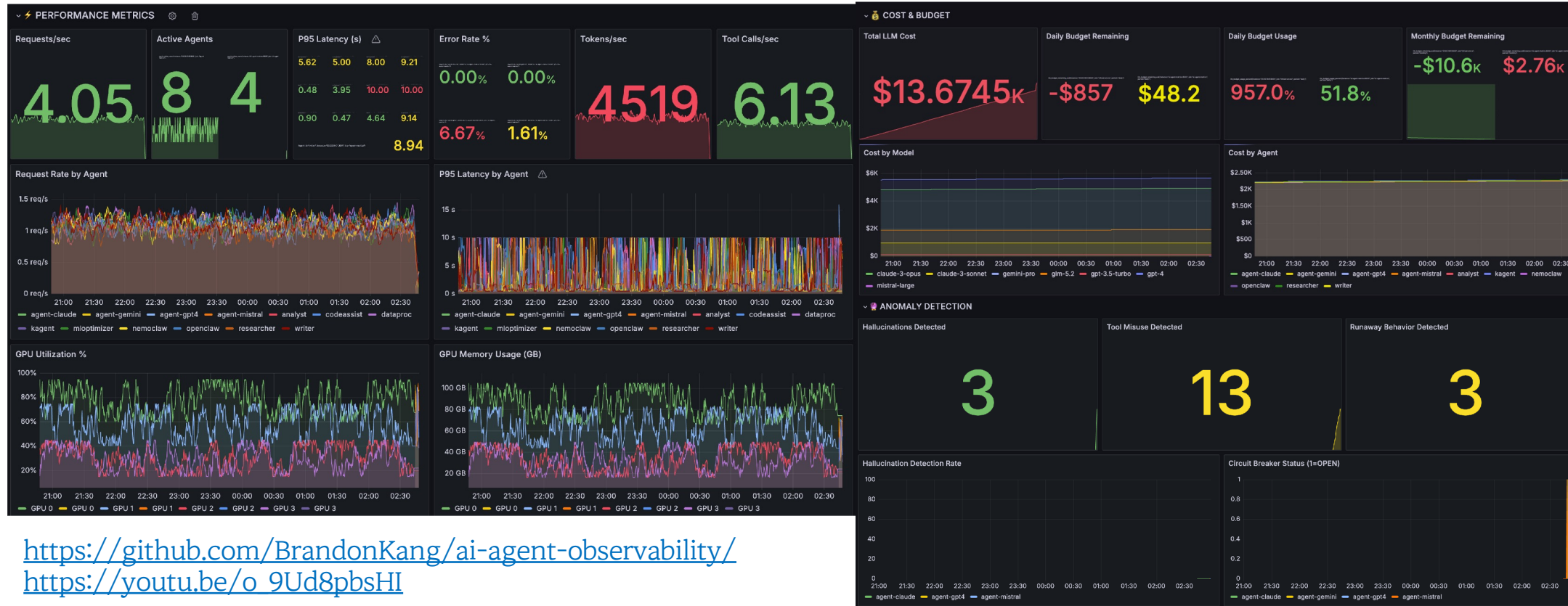


- DEMO -





Results)



<https://github.com/BrandonKang/ai-agent-observability/>
https://youtu.be/o_9Ud8pbsHI

```
1  #!/usr/bin/env python3
2  """Demo script for single agent with observability."""
3
4  import os
5  import sys
6  import time
7  import click
8  from rich.console import Console
9  from rich.panel import Panel
10
11  # Add src to path
12  sys.path.insert(0, os.path.dirname(__file__))
13
14  from src.observability.tracer import init_tracer
15  from src.observability.metrics import init_metrics
16  from src.agents.base_agent import ObservableAgent
17
18  console = Console()
19
20
21  @click.command()
22  @click.option('--task', default="Research climate change impacts in 2050", help='Task for the agent')
23  @click.option('--agent-id', default="agent-1", help='Agent ID')
24  @click.option('--model', default="GLM-5.2", help='Model name')
25  @click.option('--metrics-port', default=8000, type=int, help='Metrics port')
26  def main(task: str, agent_id: str, model: str, metrics_port: int):
27      """Run a single AI agent with observability."""
28
29      console.print(Panel.fit(
30          "[bold cyan]AI Agent Observability Demo[/bold cyan]\n"
31          "[dim]Single Agent Scenario[/dim]",
32          border_style="cyan"
33      ))
34
35  # Initialize observability
```

Impact Analysis

Before Observability:

- No cost tracking
- Runaway agent undetected
- Daily: \$847
- Monthly: \$25,400

After Implementation:

- Real-time monitoring
- Budget alerts at 70%
- Circuit breaker active
- Daily: \$127
- Monthly: \$3,800

💰 Savings: 85% (\$21,600/mo)

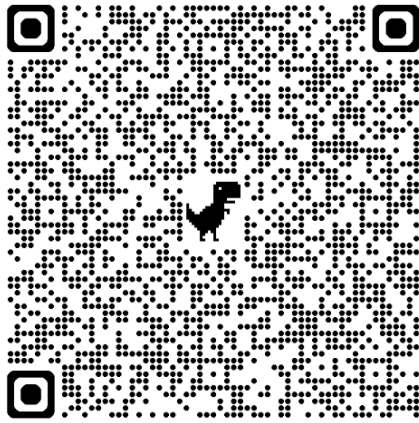
Optimization Actions:

- Simple tasks → GLM-4
 - (cheaper than GLM-5.2)
- Batch processing
 - NemoClaw: 50→200 tasks
- Model selection by task
 - Classification → GLM-4
 - Generation → GLM-5.2
- Result: 40% further reduction



- Q&A -





Slide Download



Demo GitHub Repo



Key Takeaways

- Observability is NOT optional for AI agents
→ Silent failures are common without visibility
- Multi-layered anomaly detection is essential
→ 4 detectors → 94% precision vs single method ~60%
- Cost tracking drives optimization
→ What gets measured gets managed - 85% cost reduction
- Open source scales to production
→ Prometheus + Grafana: 1M+ metrics/day, zero lock-in



서울

Thank You!!

Connect with us:

 Brandon Kang sakang@akamai.com

 Mostafa Radwan mr@datadog.com

