



Sky Computing for Serverless: Infrastructure Assessment to Support Performance Enhancement

Robert Cordingly, Xinghan Chen, Ling-Hong Hung, Wes Lloyd
rcording@uw.edu

School of Engineering and Technology
University of Washington Tacoma

IEEE/ACM 18th International Conference on Utility and Cloud
Computing (UCC 25) December 1-4, 2025, Nantes France



Background and Motivation



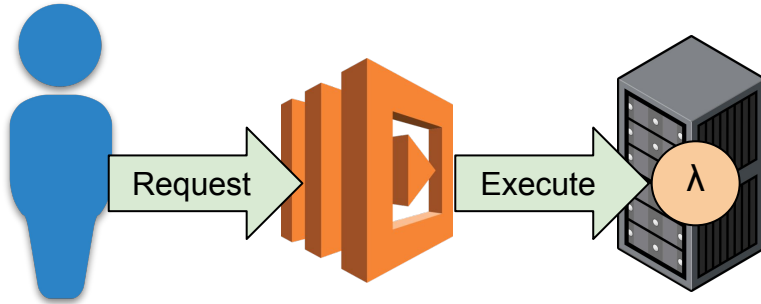
What is Serverless?



Serverless function-as-a-service (FaaS) platforms offer many desirable features:

- Rapid elastic scaling
- Scale to zero
- No infrastructure management
- Fine grained billing
- Fault tolerance
- No up front cost to deploy an application

How FaaS Platforms Work



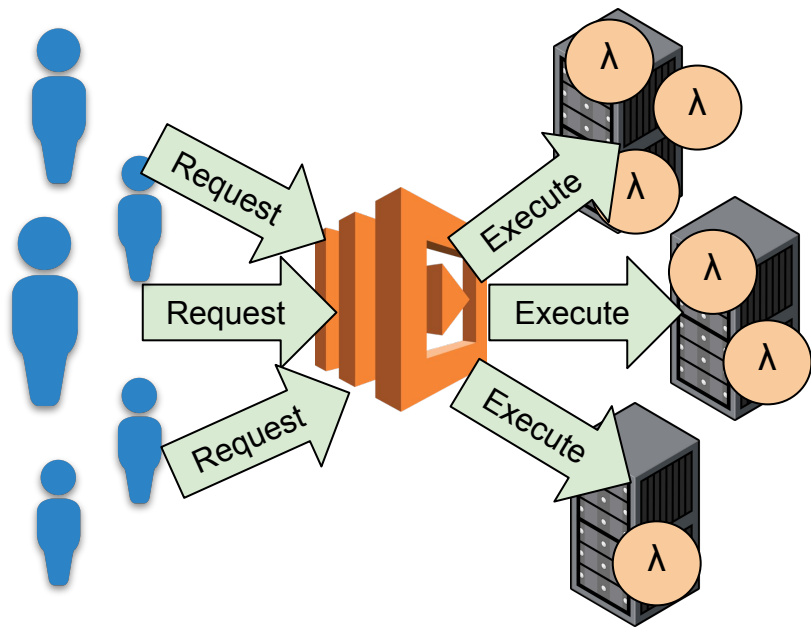
Client

FaaS
Platform

Function
Instance

- Clients make requests to the FaaS platform which manage the infrastructure automatically
- FaaS platforms create and execute users code inside environments known as function instances
- Function instances can be hosted on a variety of different kinds of hardware

Automatic Scaling



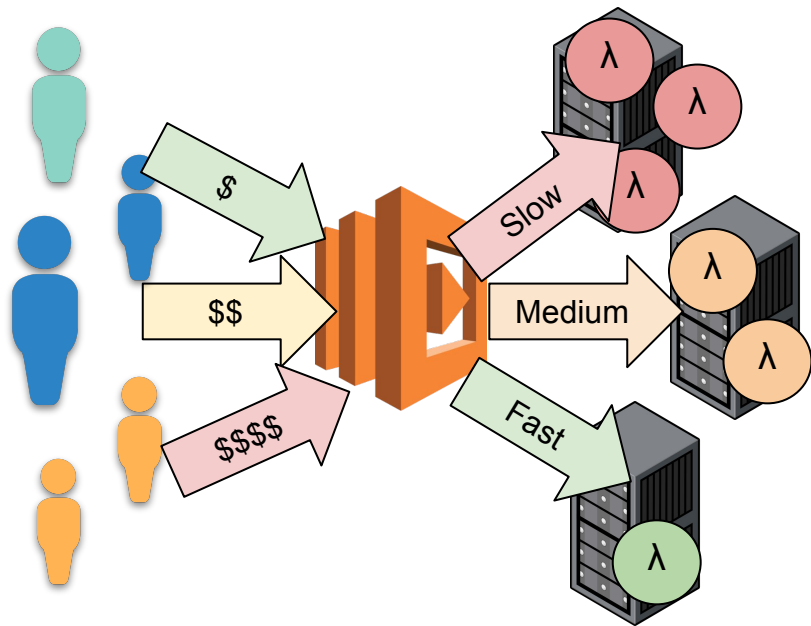
Clients

FaaS
Function

Function
Instances

- As more requests are made, the FaaS platform scales the number of function instances
- Function instances for one function can be spread across many host servers

Inconsistent Performance



Clients

FaaS
Function

Function
Instances

- Heterogeneous hardware leads to some hosts being faster than others
- Since FaaS platforms bill based off runtime, slower hosts perform worse and are more expensive

Request Routing



- If we were able to observe all of the available infrastructure of serverless regions, we could route requests to regions with more favorable infrastructure
- Our previous work built a serverless sky proxy system to route requests around the world (IC2E 2023)

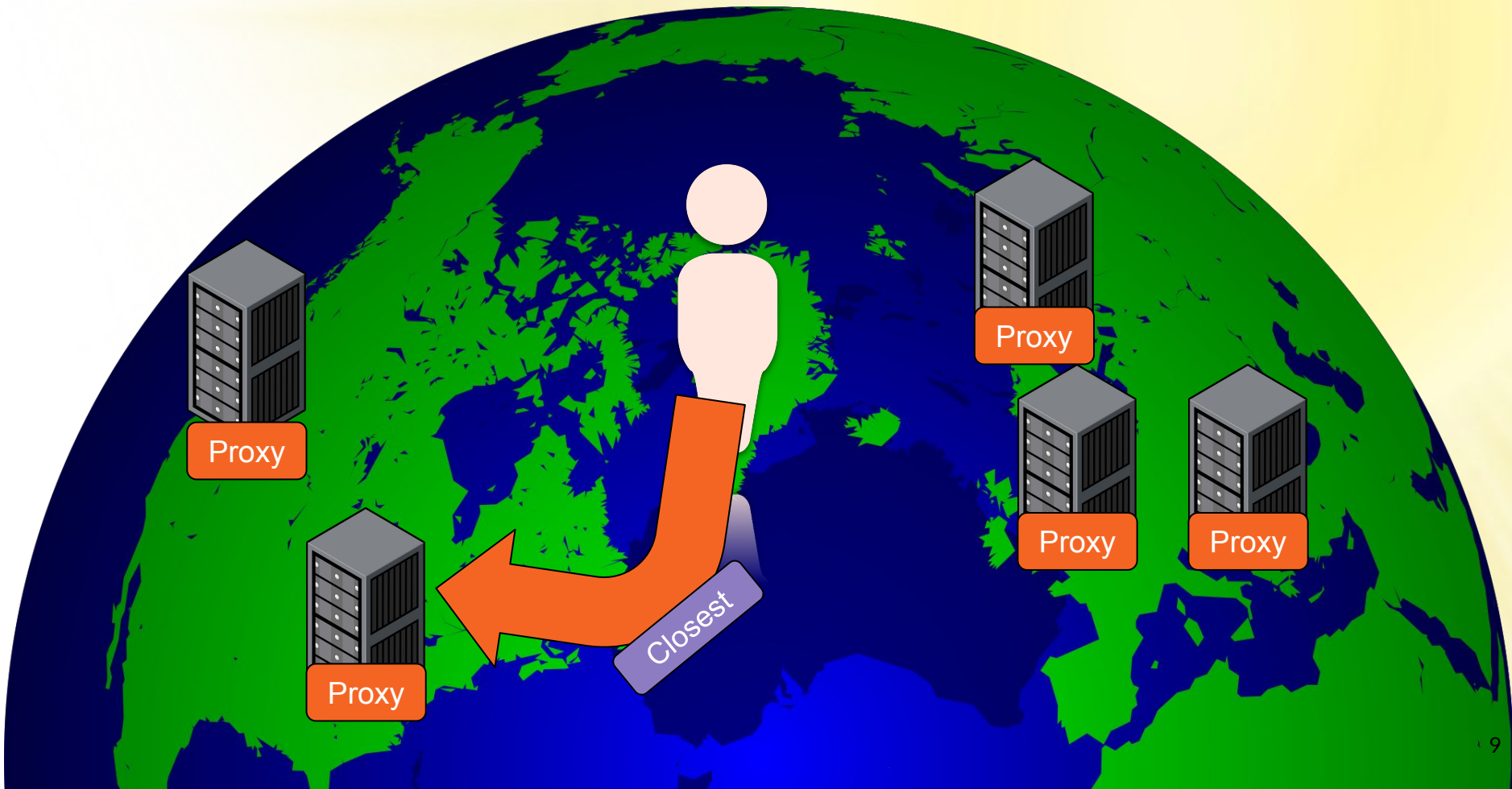
What is Sky Computing?

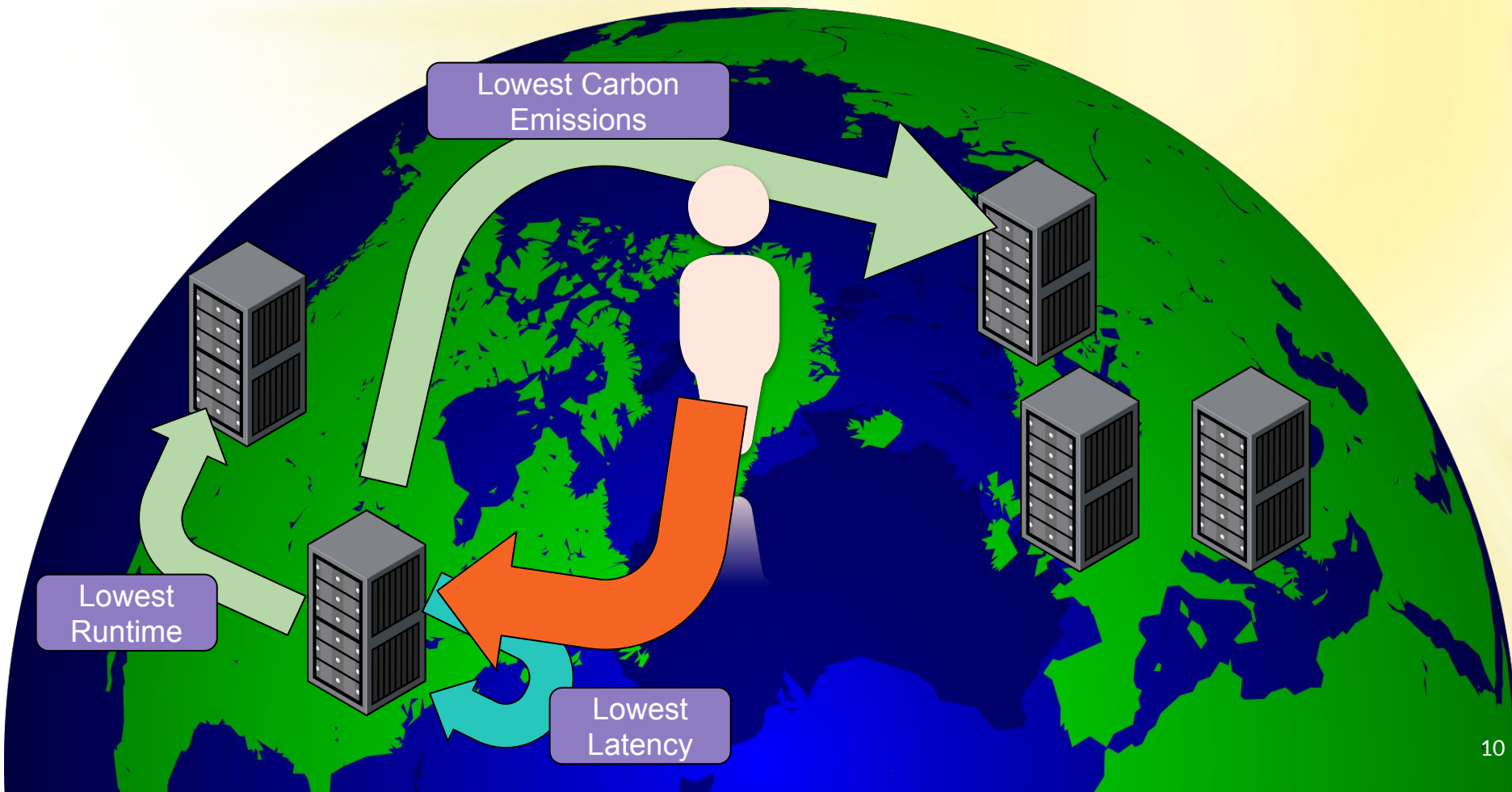


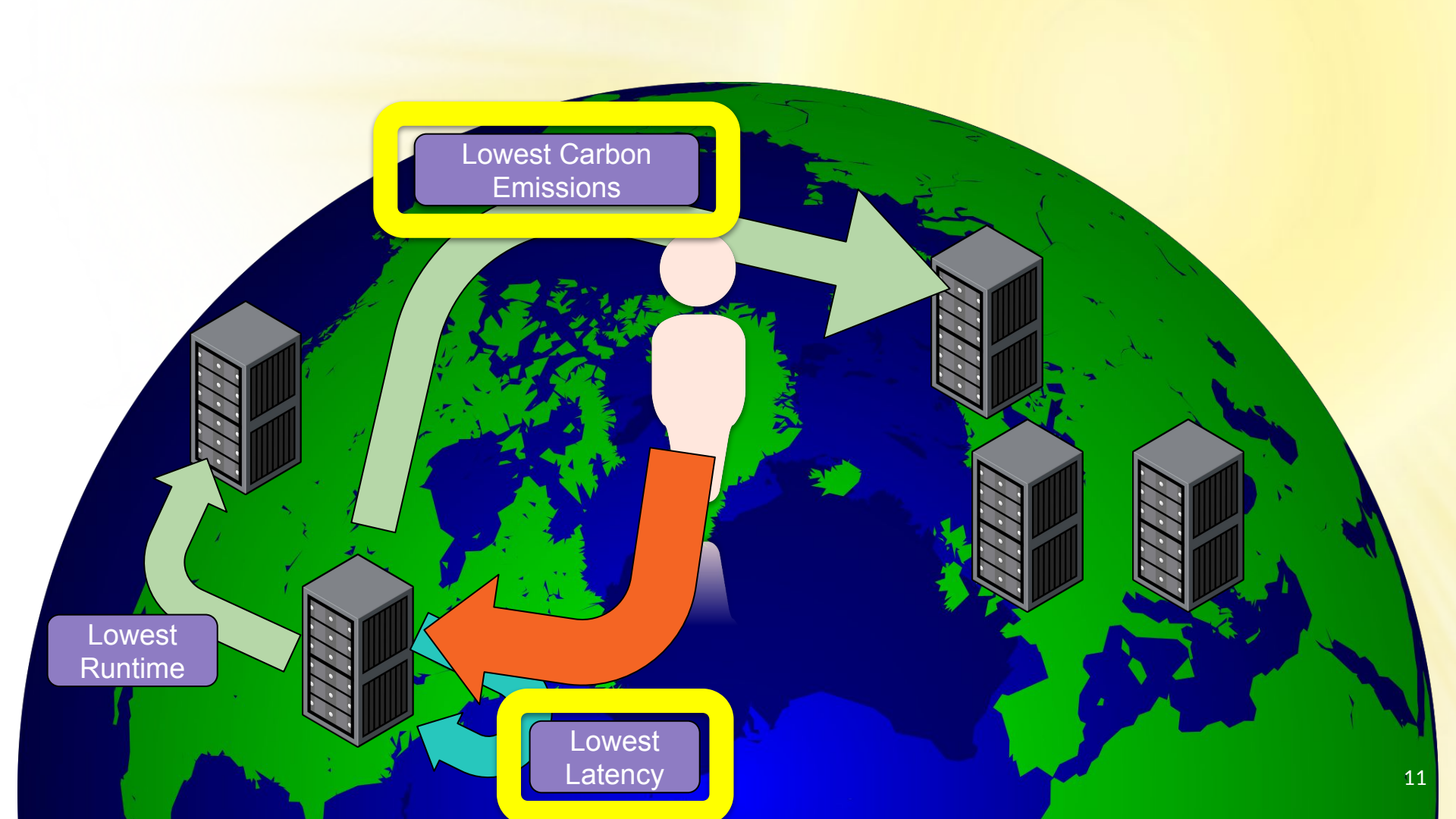
- The Sky sits above the clouds
- Consists of compatibility layers allowing aggregation of resources between cloud regions, availability zones, or cloud providers

Goals for Serverless Sky Computing:

- Allow applications to take advantage of resources of multiple regions and cloud providers
- Improve performance and reduce costs of FaaS applications





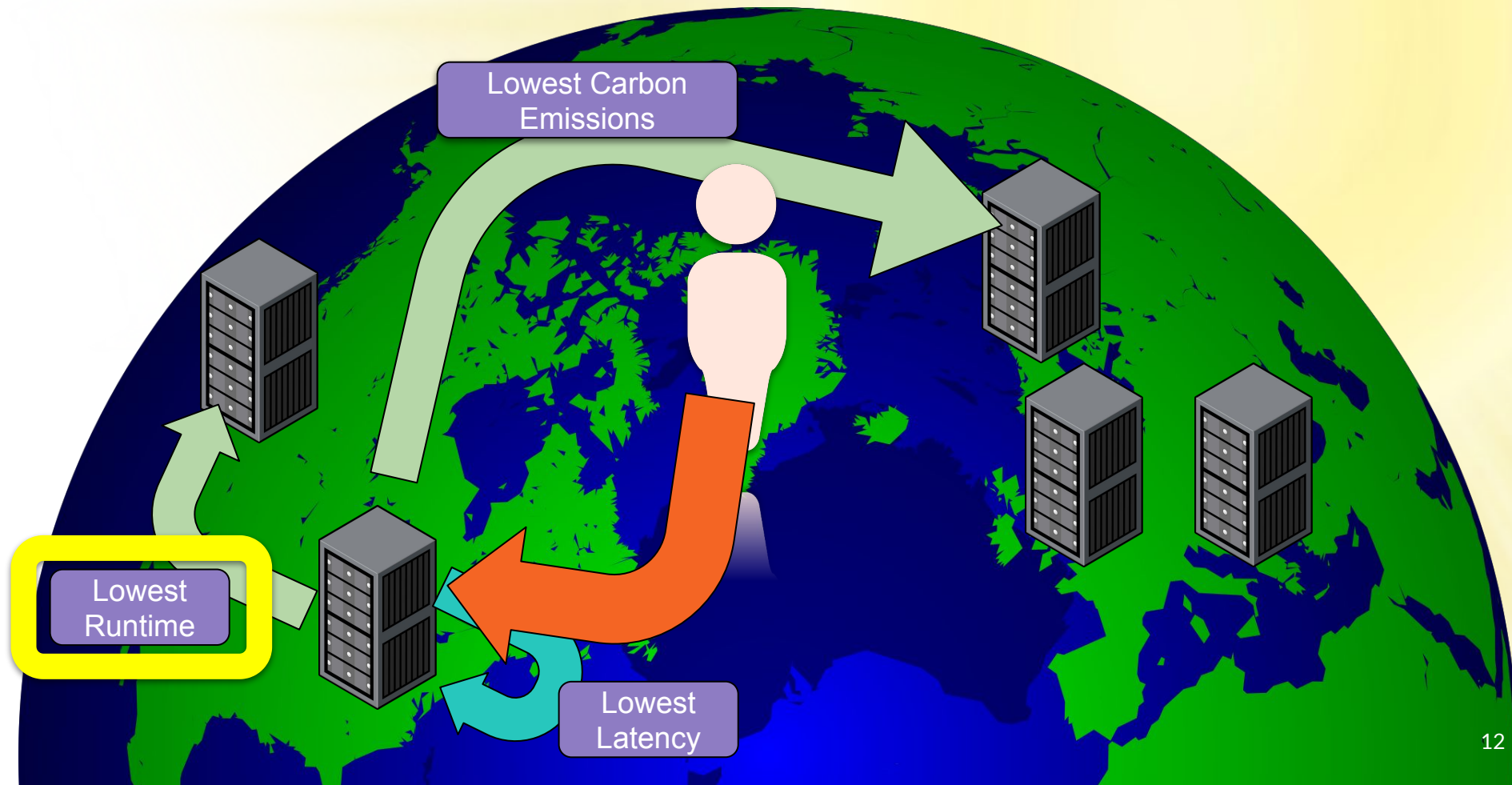


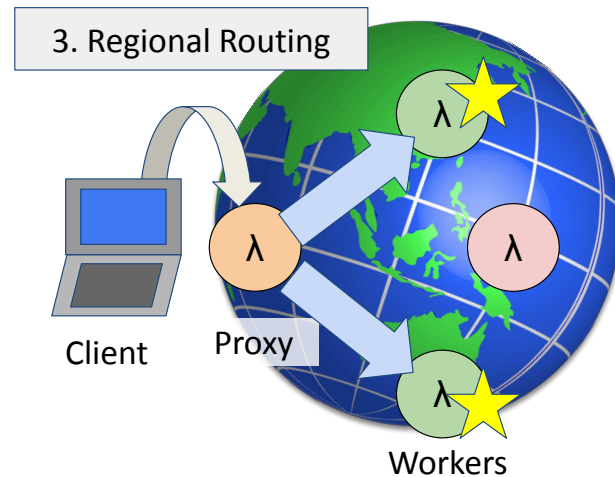
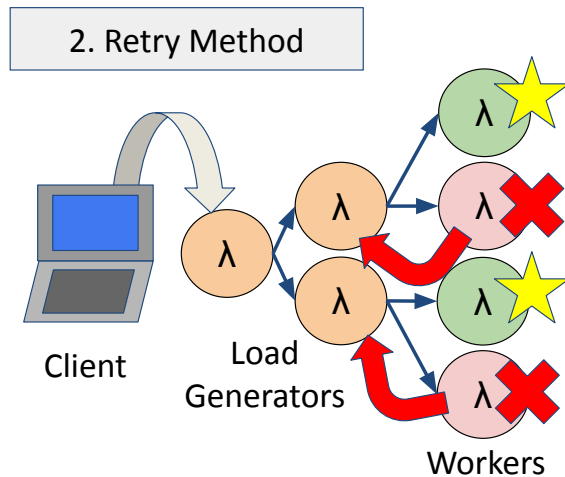
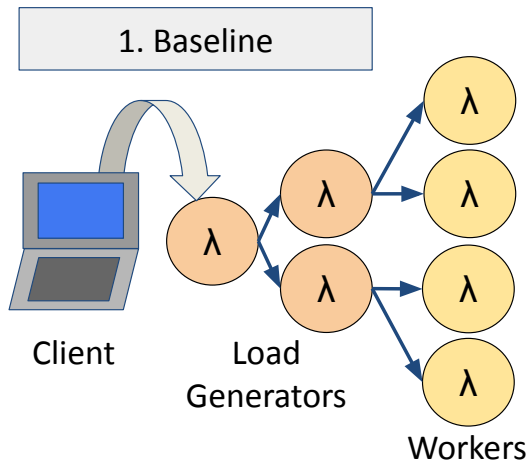
The diagram shows a stylized Earth with a person standing in the center. Five server racks are distributed across the globe. Three arrows originate from the person: a light green arrow pointing to a server in North America, a light green arrow pointing to a server in Europe, and a thick orange arrow pointing to a server in Asia. Three labels are present: 'Lowest Carbon Emissions' in a yellow box at the top, 'Lowest Runtime' in a purple box on the left, and 'Lowest Latency' in a yellow box at the bottom.

Lowest Carbon Emissions

Lowest Runtime

Lowest Latency





Retry Logic

1. **Retry Slow:** Invocations retry on the two slowest CPUs
2. **Focus Fastest:** Invocations retry on everything but the fastest CPU

Hybrid Approach

Combine both Retries and Regional Routing

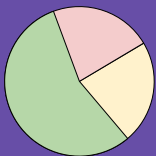


Research Questions

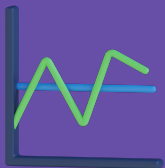
Research Questions



- **RQ-1 (Infrastructure Variation):** What CPU variations can be observed on FaaS platforms across different zones and regions? How does it change over time?



- **RQ-2 (Infrastructure Characterization):** How many samples are required to accurately infer the hardware pool? How can we balance the cost versus accuracy?



- **RQ-3 (Performance Optimization):** To what extent are runtime and cost improvements possible by targeting specific instructure?

Serverless Infrastructure Observation

Using SAAF in a Function:

Using SAAF in a function is as simple as importing the framework and adding a couple lines of code. Attributes collected by SAAF will be appended onto the JSON response. For asynchronous functions, this data could be stored into a database, such as AWS S3, and retrieved after the function is finished.

Example Function:

```
from Inspector import *

def myFunction(request):
    # Initialize the Inspector and collect data.
    Inspector = Inspector()
    Inspector.inspectAll()

    # Add a "Hello World!" message.
    Inspector.addAttribute("message", "Hello " + request['name'] + "!")

    # Return attributes collected.
    return Inspector.finish()
```

Example Output JSON:

The attributes collect can be customized by changing which functions are called. For more detailed descriptions of each variable and the functions that collect them, please see the framework documentation for each language.

```
{
  "version": "0.3",
  "lang": "python",
  "cpuType": "Intel(R) Xeon(R) Processor @ 2.58GHz",
  "cpuModel": "63",
  "vmuptime": "1551727835",
  "uuid": "6241c618-7808-48e2-8736-997dc1a9316d",
  "vendor": "ELUCOA",
  "platform": "AWS Lambda",
  "newContainer": "1",
  "cpuIdleDelta": "984",
  "cpuUsedDelta": "0",
  "cpuIdleDelta": "385",
  "cpuUsedDelta": "82428",
  "cpuIdleDelta": "226",
  "cpuUsedDelta": "0",
  "cpuIdleDelta": "0",
  "cpuUsedDelta": "1594",
  "frameworkUptime": "15,12",
  "message": "Hello Fred Smith!",
  "runtime": "38.94"
}
```

Attributes Collected by Each Function

The amount of data collected is determined by which functions are called. If some attributes are not needed, then some functions may not need to be called. If you would like to collect every attribute, the `inspectAll()` method will run all methods.

Core Attributes	
Field	Description
version	The version of the SAAF Framework.
lang	The language of the function.
runtime	The server-side runtime from when the function is initiated until <code>Inspector.finish()</code> is called.
startTime	The Unix Epoch that the Inspector was initialized in ms.

InspectContainer()	
Field	Description
uuid	A unique identifier assigned to a container if one does not already exist.
newContainer	Whether a container is new (no assigned uuid) or if it has been used before.
vmuptime	Time when the host booted in seconds since January 1, 1970 (Unix epoch).

InspectCPU()	
Field	Description
cpuType	The model name of the CPU.
cpuModel	The model number of the CPU.

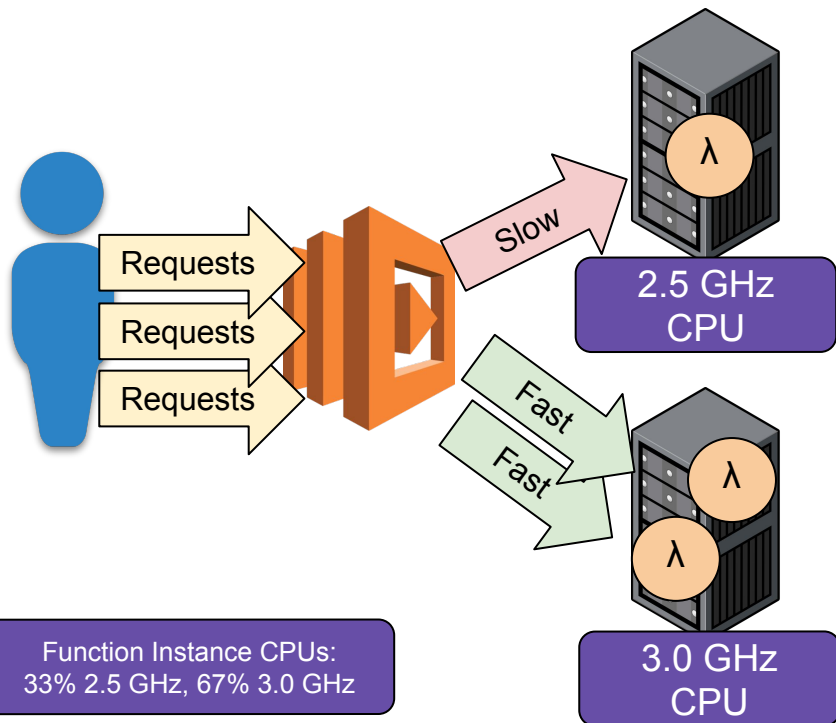
Supporting Tools - SAAF

We utilize the **Serverless Application Analytics Framework** to collect various metrics about the infrastructure and platform serverless functions are hosted with.

SAAF collects CPU Timing metrics, latency information, hardware specifications, runtime metrics, and more.

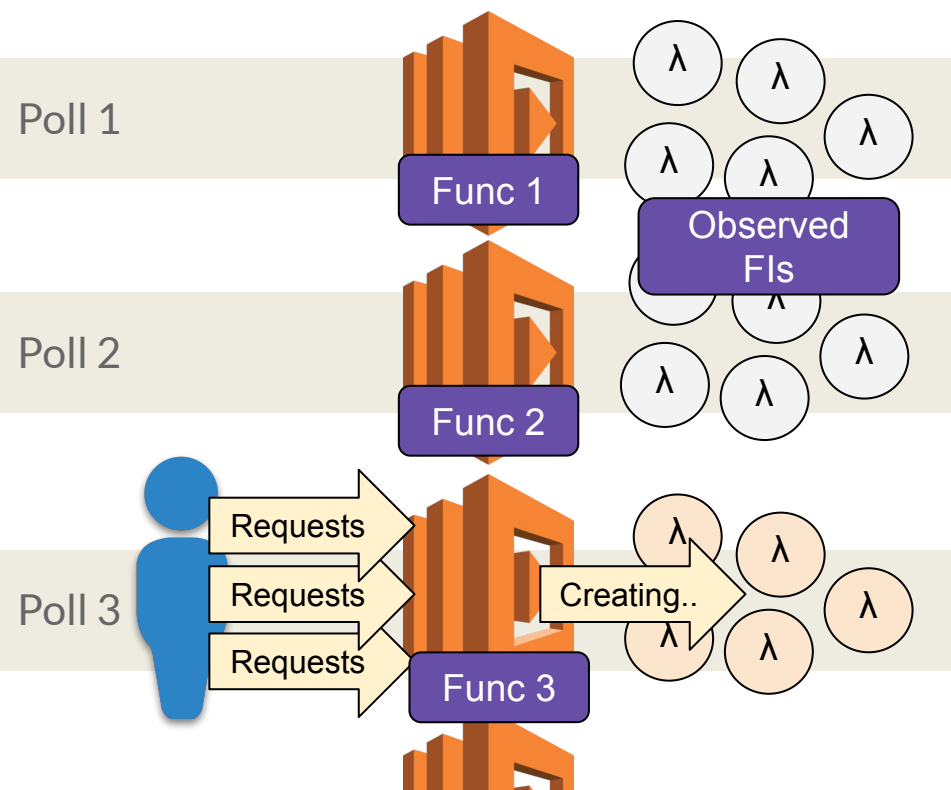
We utilize CPU metrics from SAAF to characterize function instances and observe details about the host infrastructure.

Observing Infrastructure

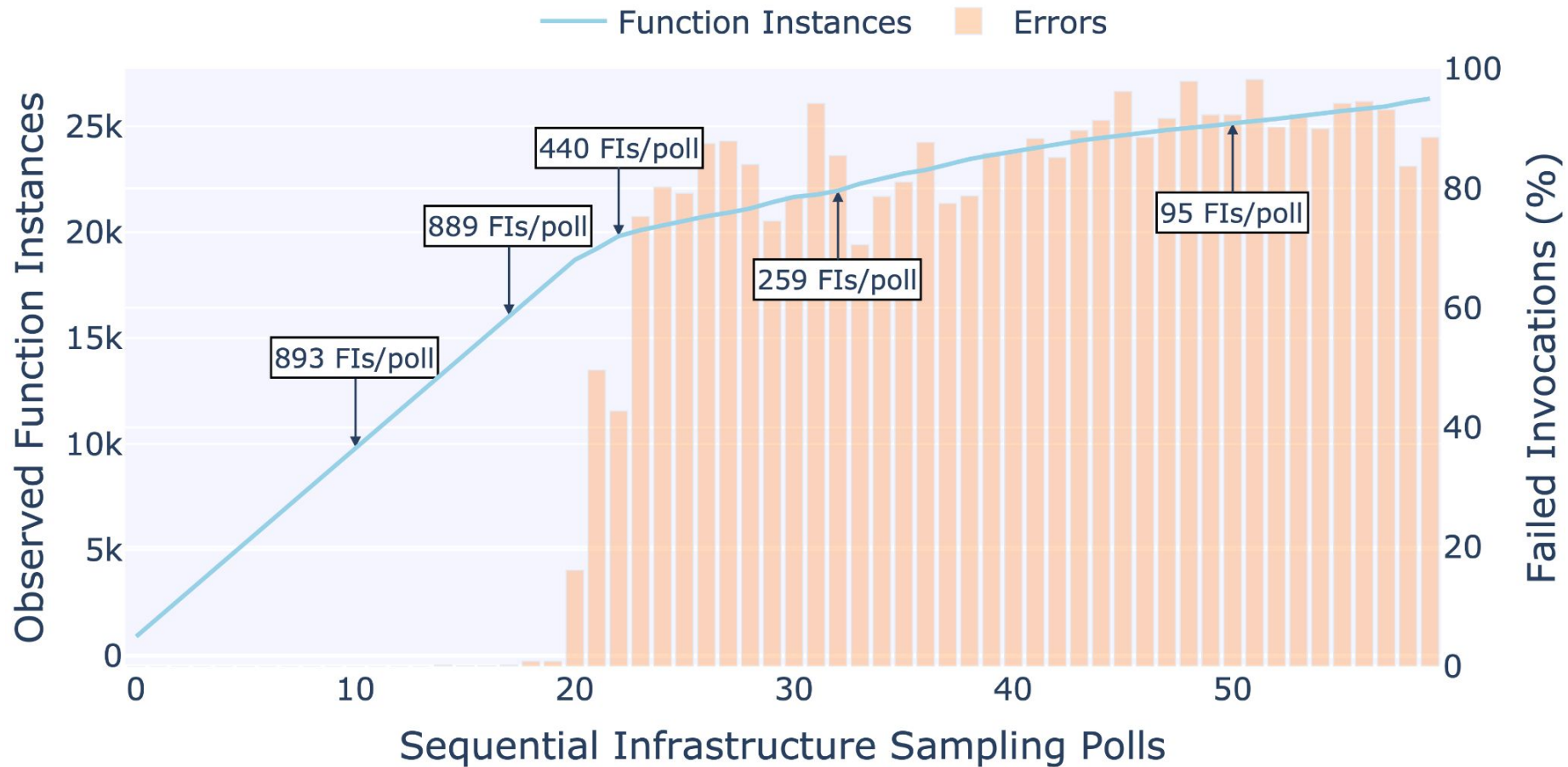


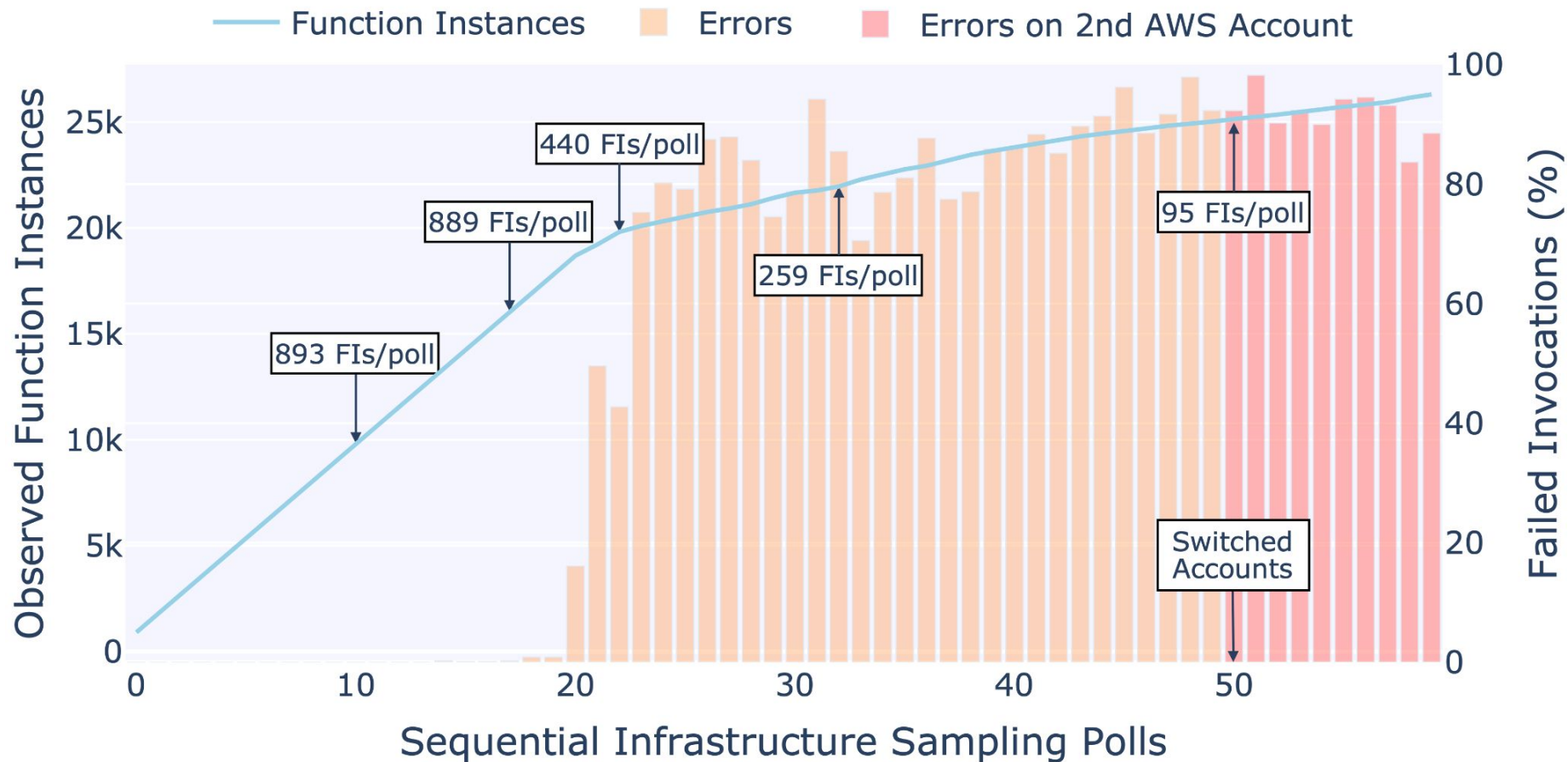
- We deployed “Hello World” sleep functions with SAAF that collect CPU information
- We then make parallel requests to view available infrastructure for a function
- With SAAF’s data we build a characterization of the pool of available infrastructure
- But this method is limited...

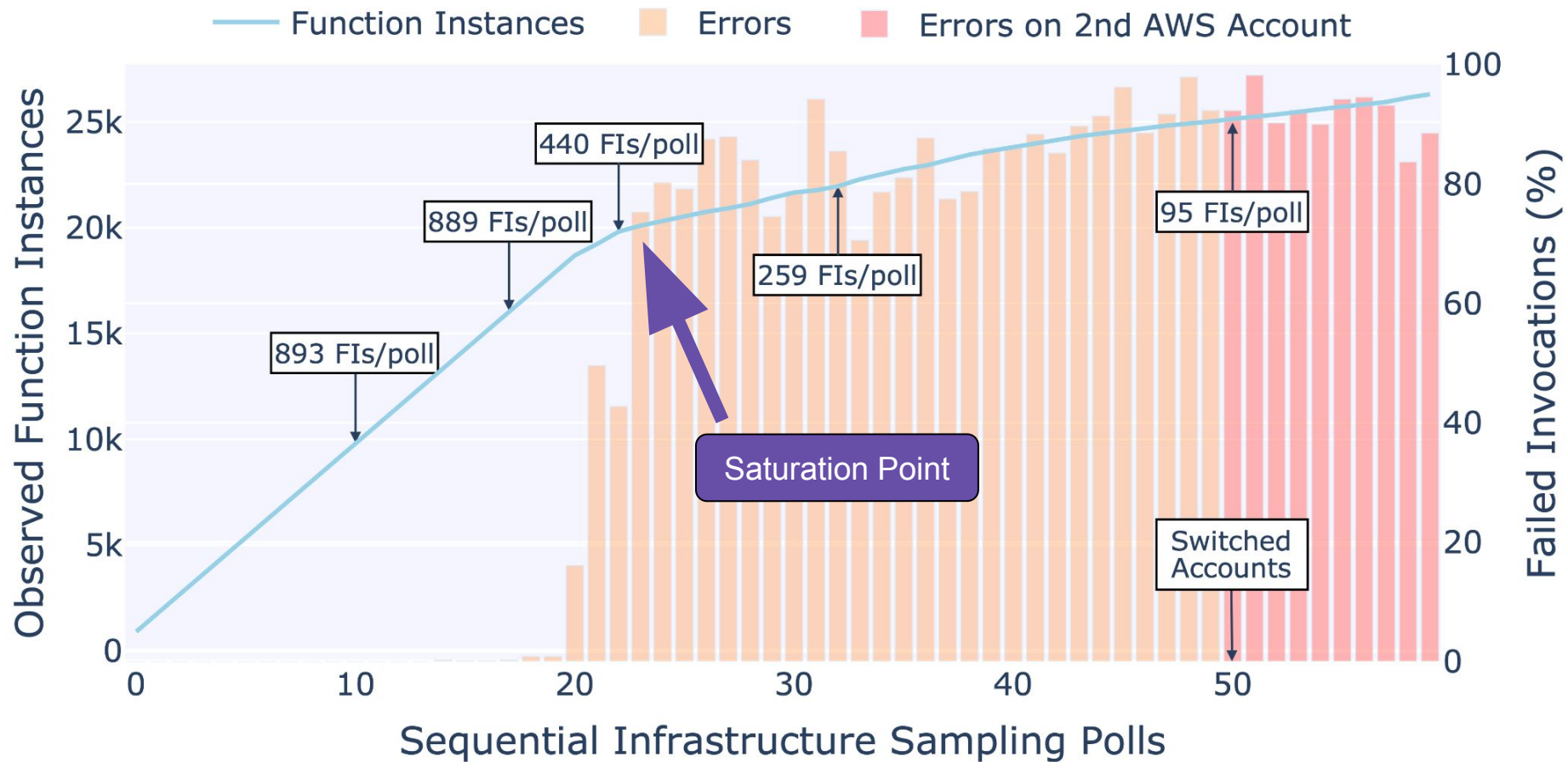
Creating Function Instances



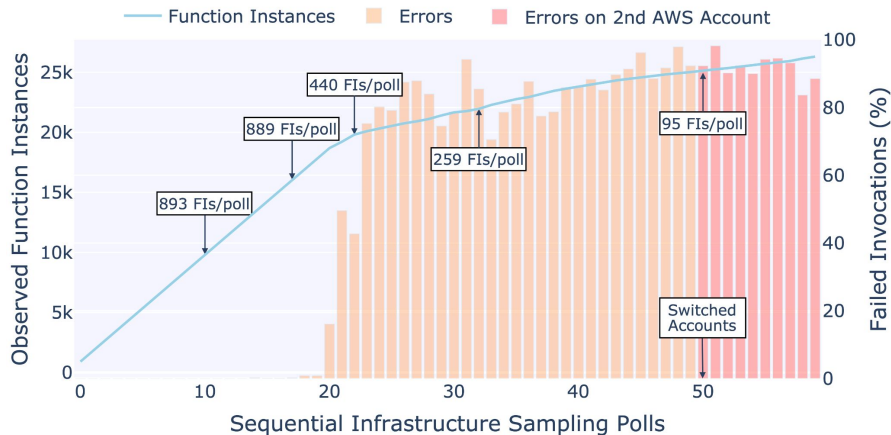
- To observe as much infrastructure as possible we need to create many function instances
- FaaS platforms limit the number of concurrent executions (1000 on AWS Lambda)
- By utilizing many function deployments, we can create more function instances than the concurrency limit
- With this we can build accurate characterizations of available CPUs in availability zones (AZs)







Sanity Checks

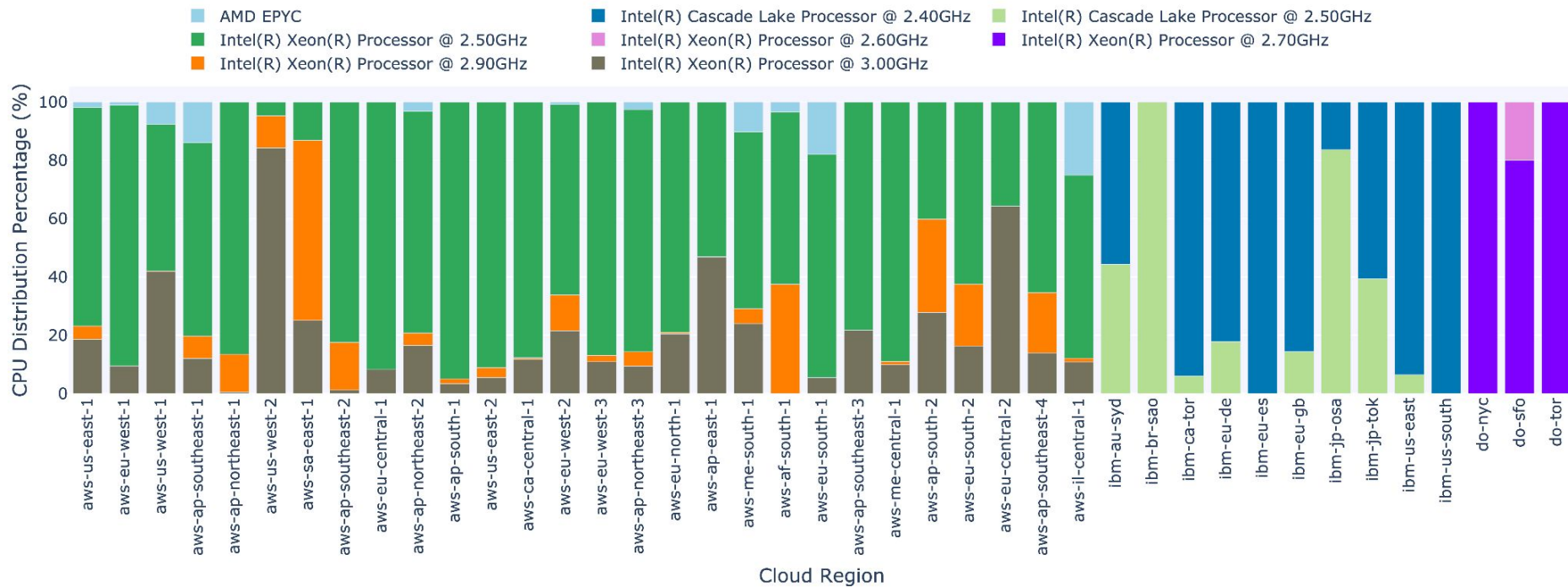


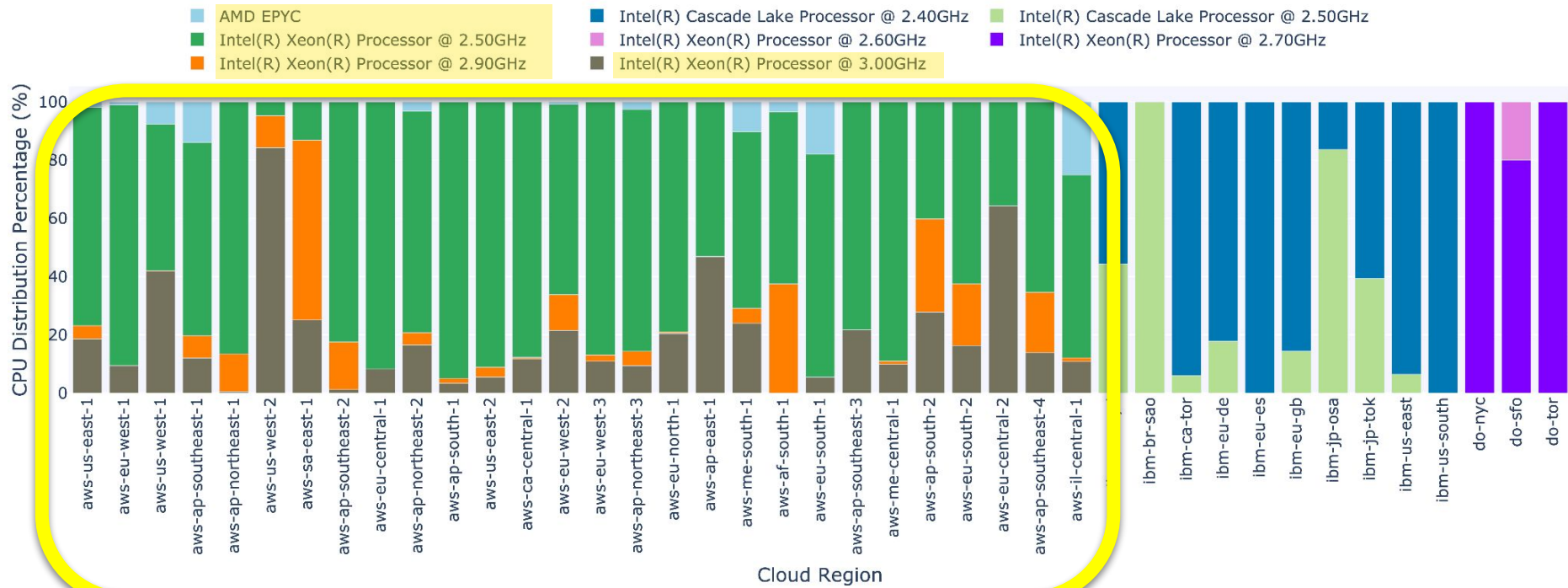
- To verify this was not a simple rate limit we:
 - used separate AWS accounts with:
 - different email addresses
 - different payment methods
 - different locations/network
 - different deployed functions
 - different FaaS configurations (checksum/RAM)
 - different invocation methods:
 - Function URLs
 - API Gateway
 - AWS CLI
- The only similarity was functions shared an availability zone
- Even with all of these tests we still observed these errors

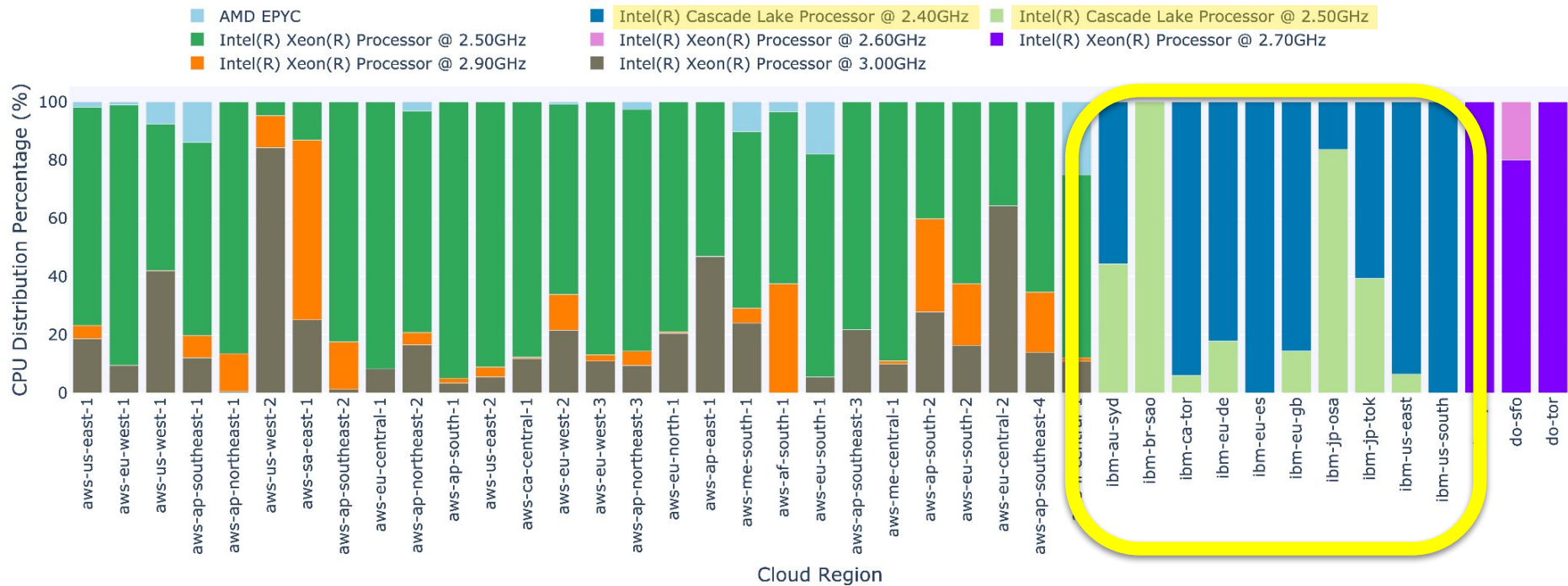


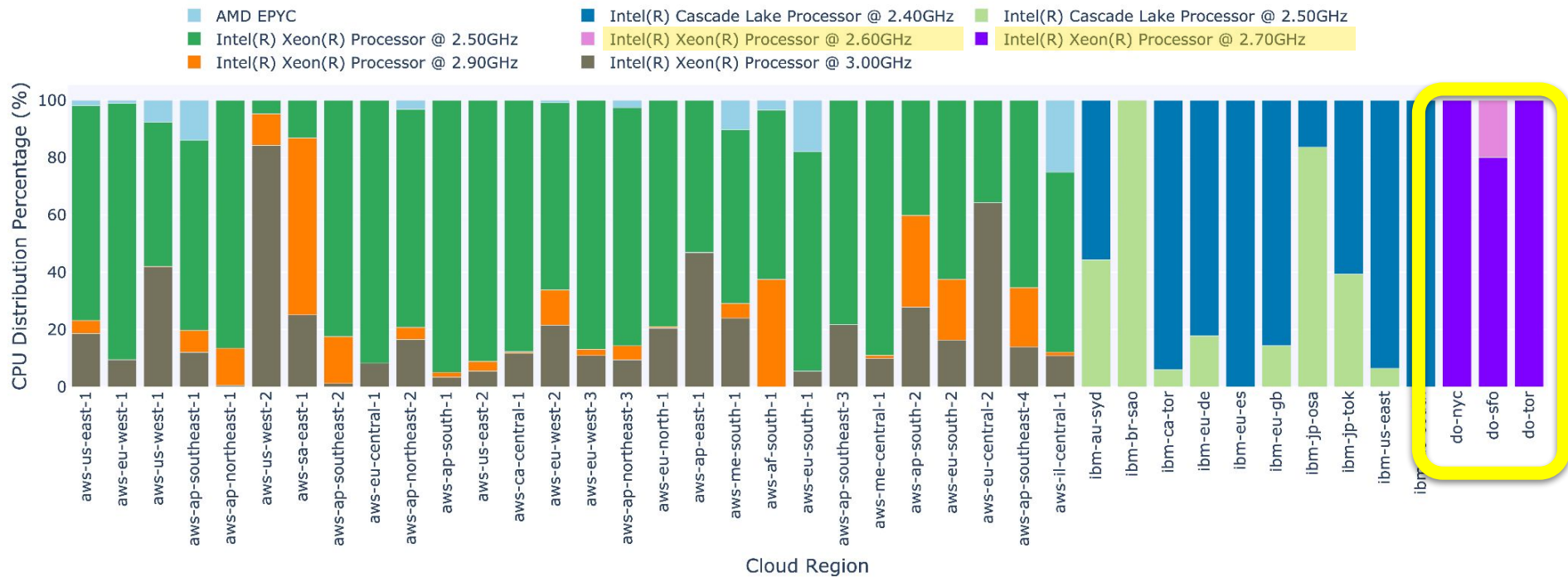
Methodology and Results

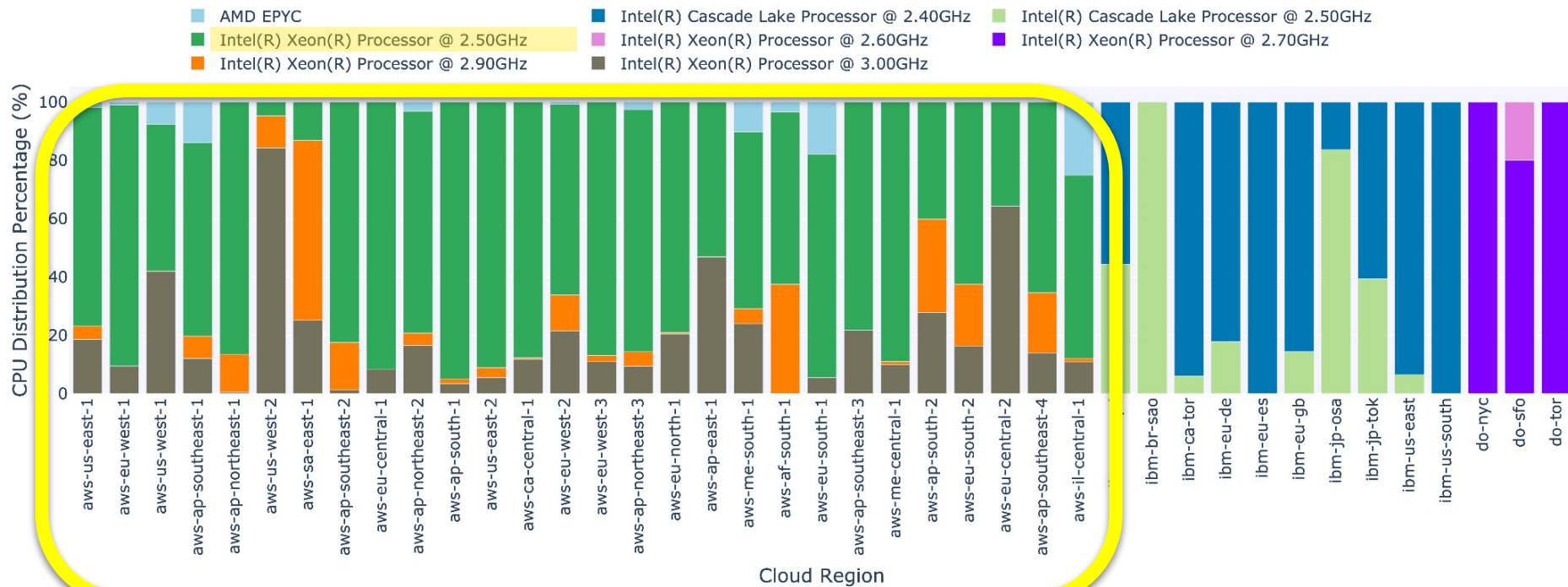
**RQ-1 (Part 1): What CPU
variations can be observed on
FaaS platforms across
different zones and regions?**

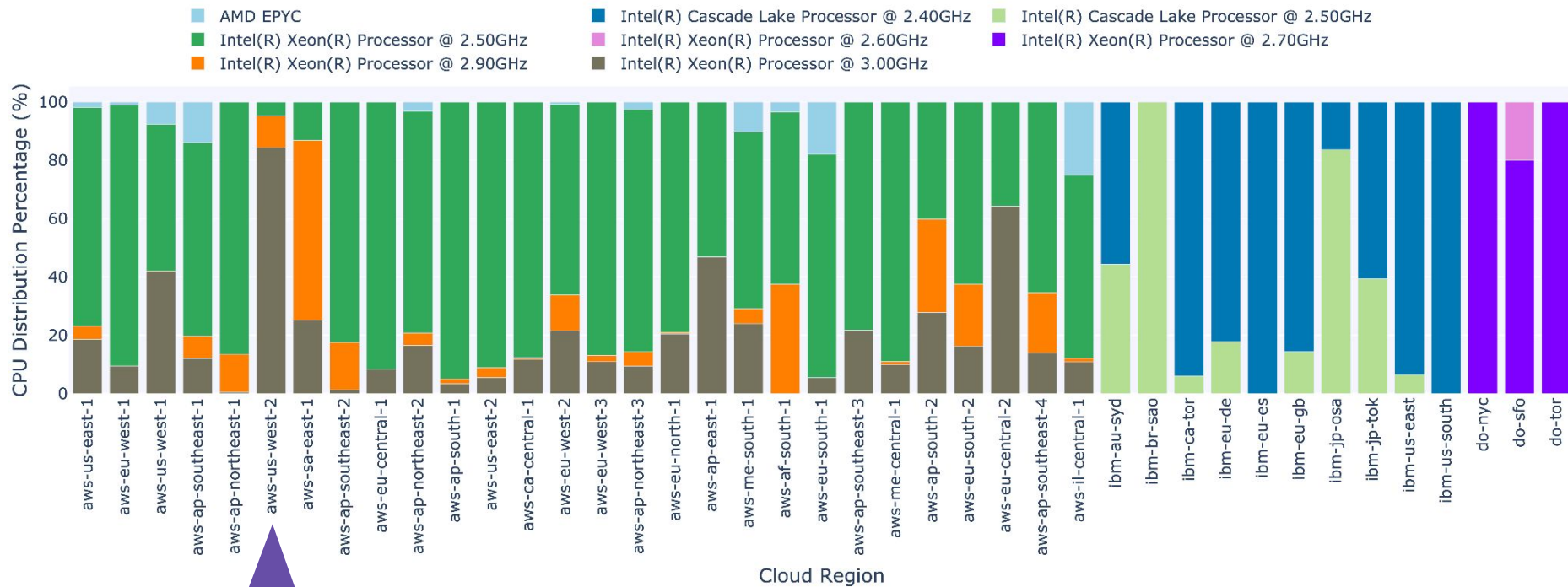


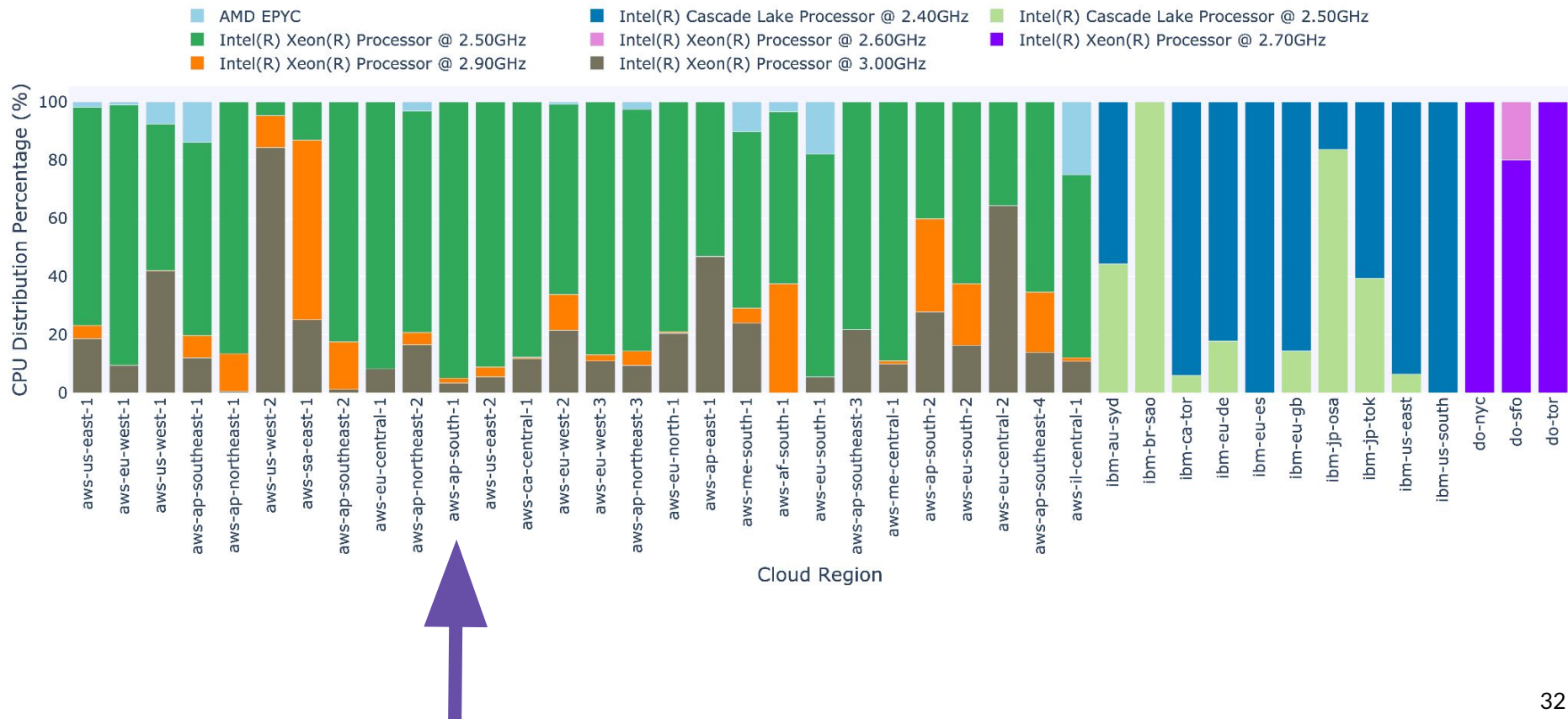


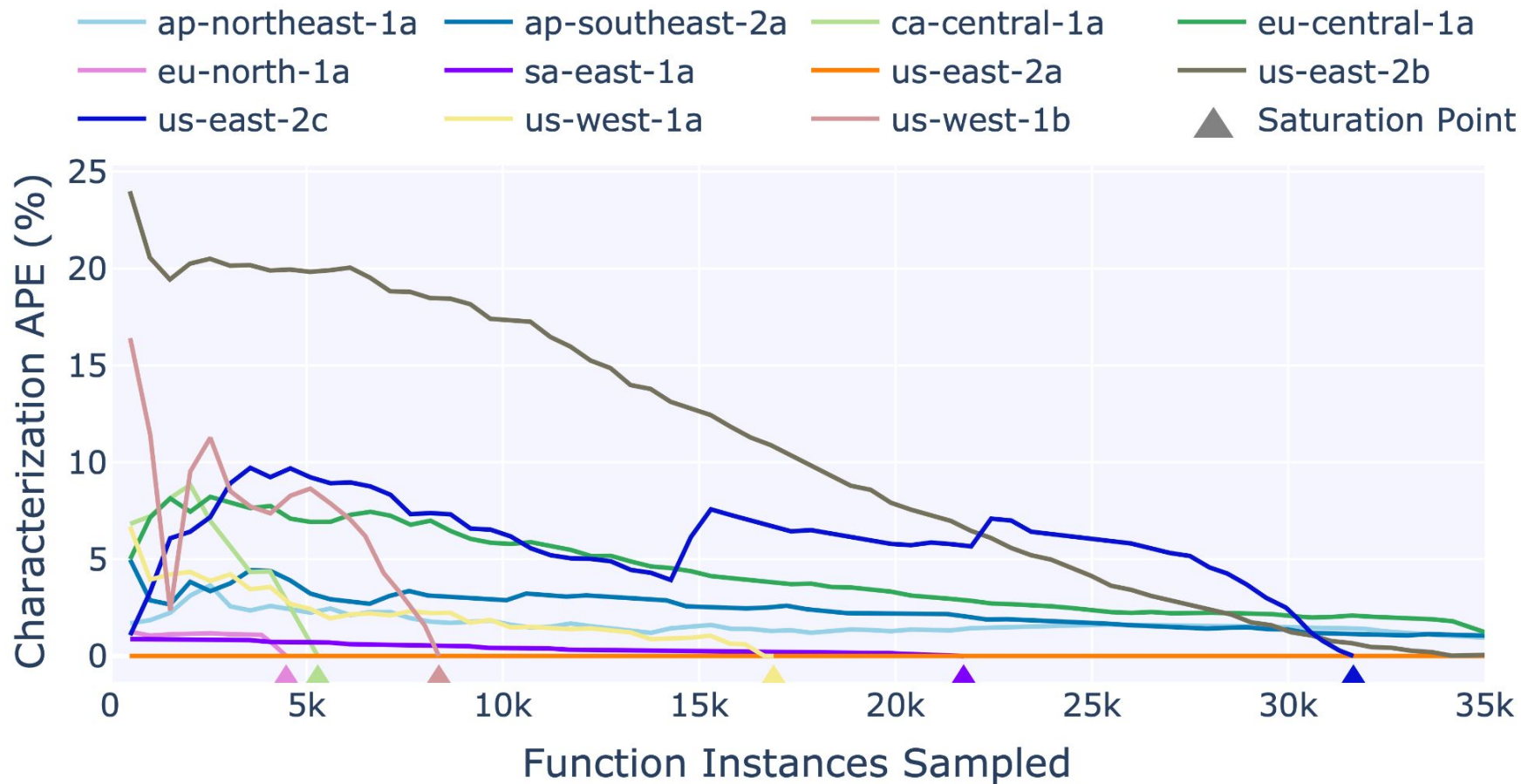


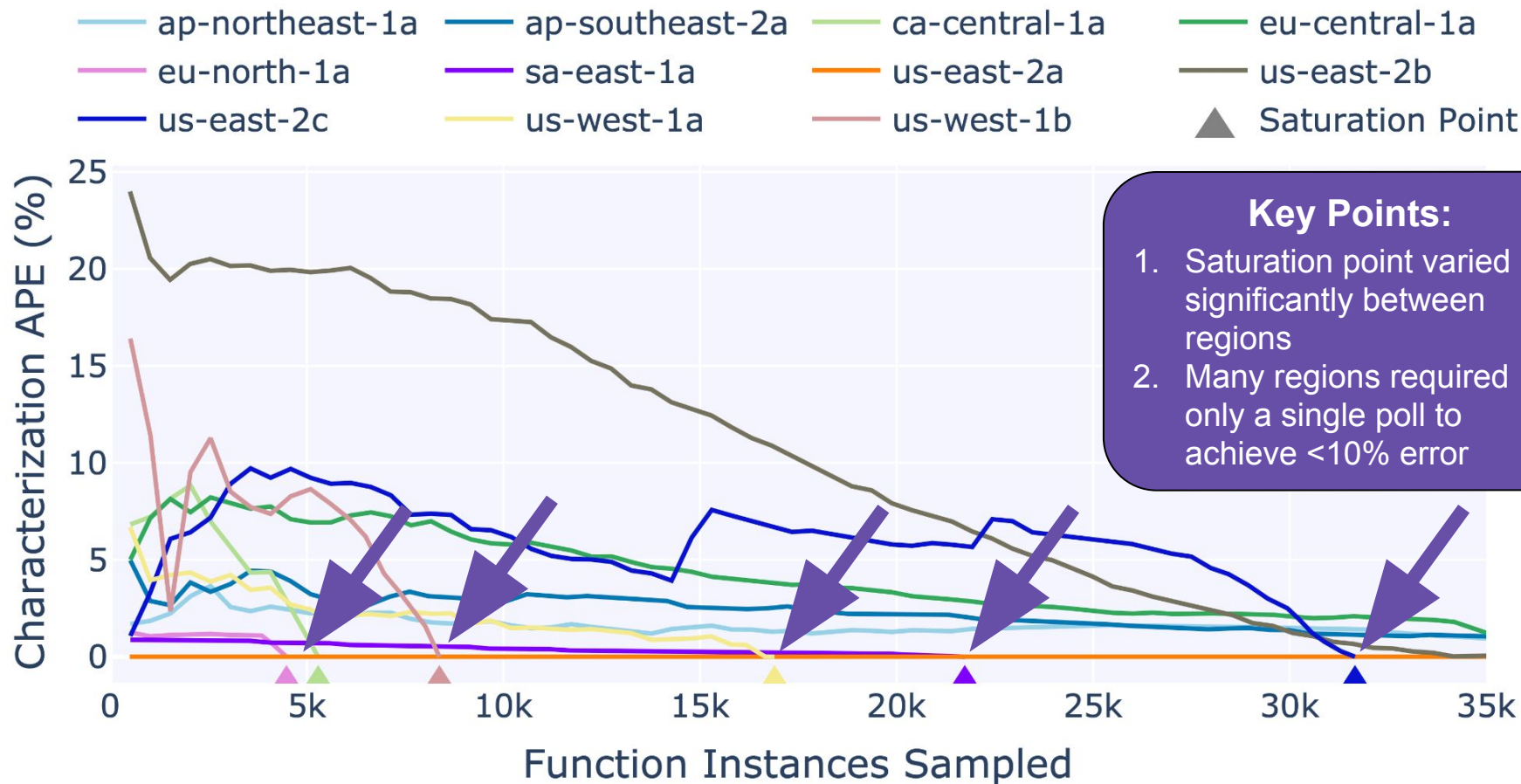


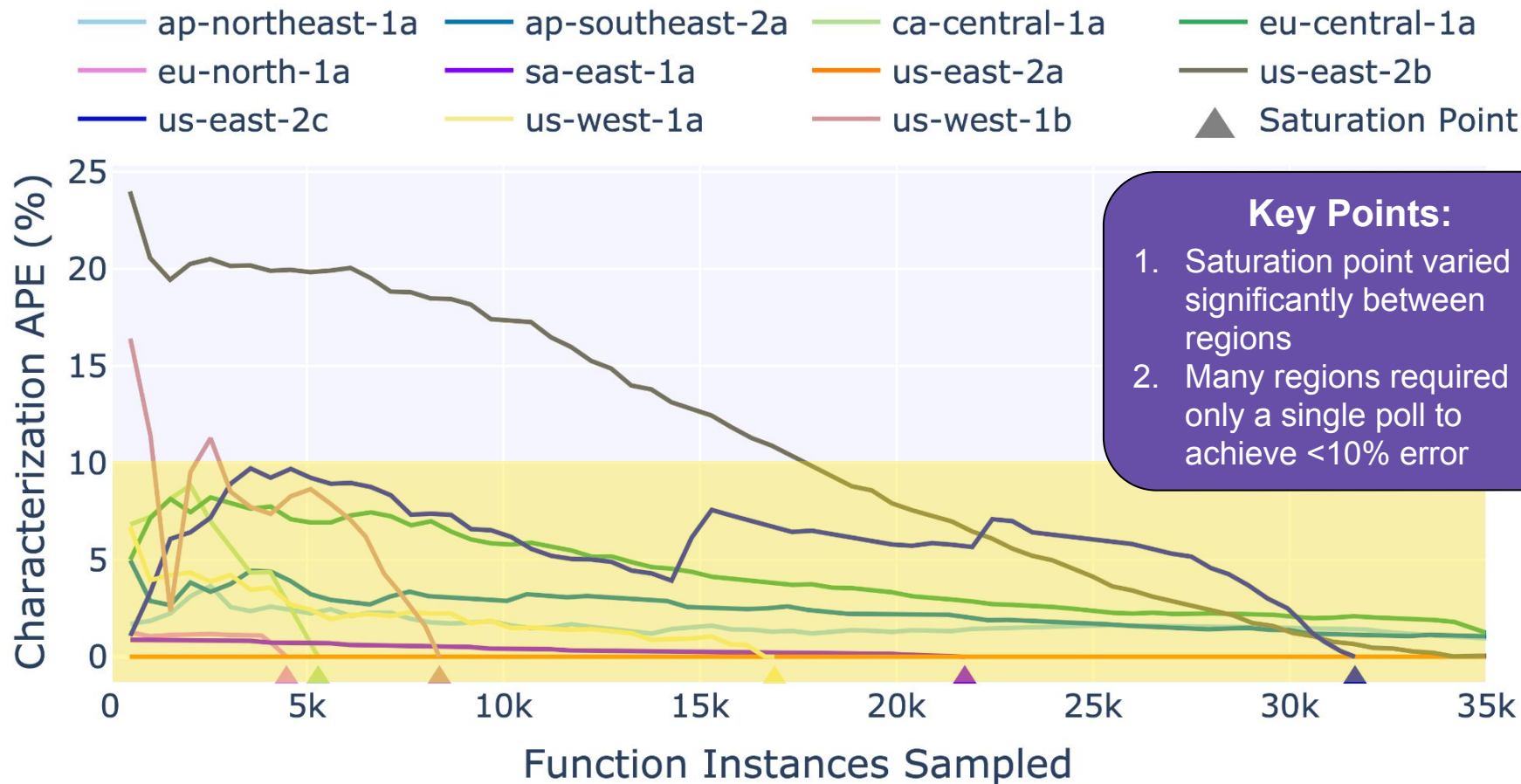


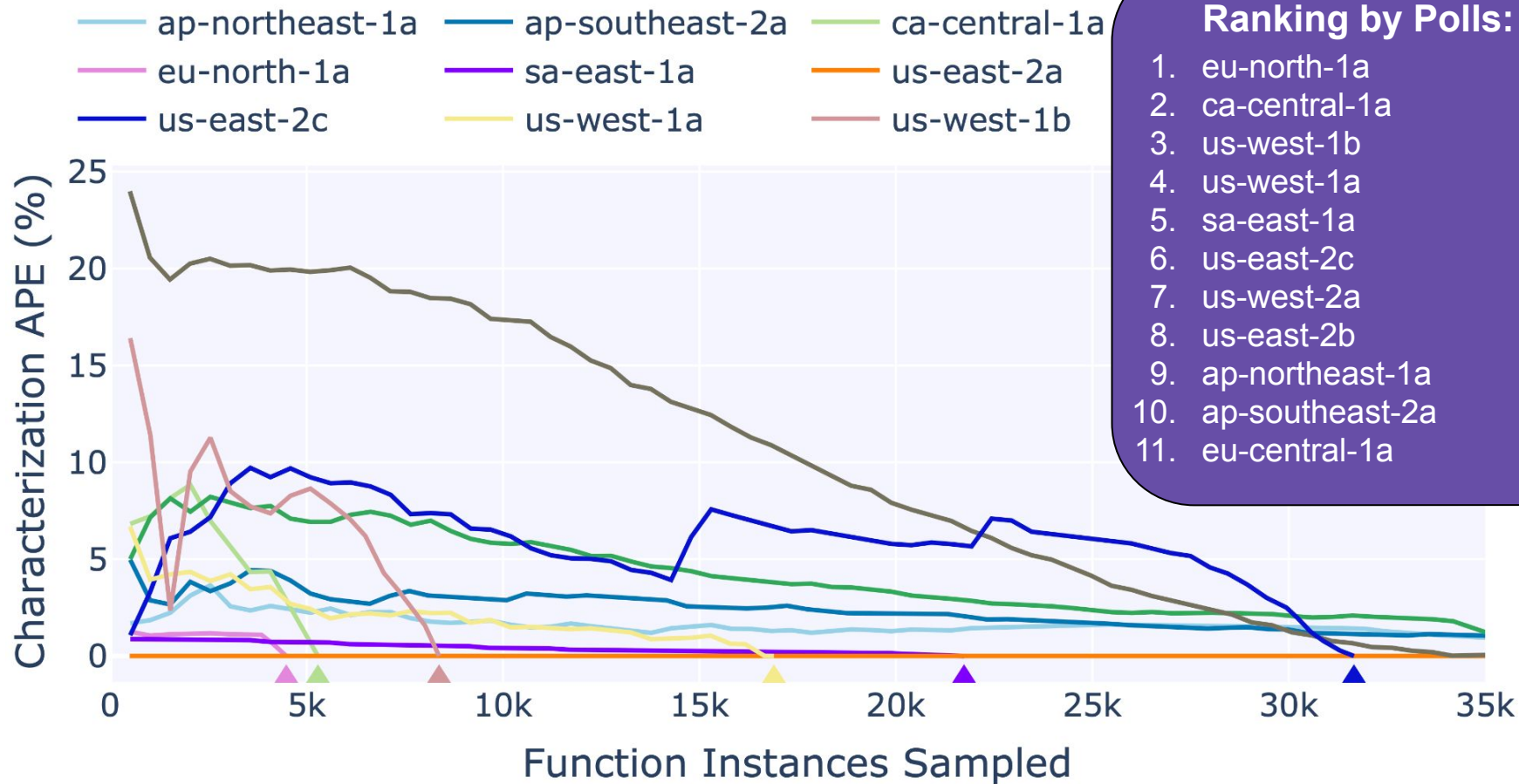




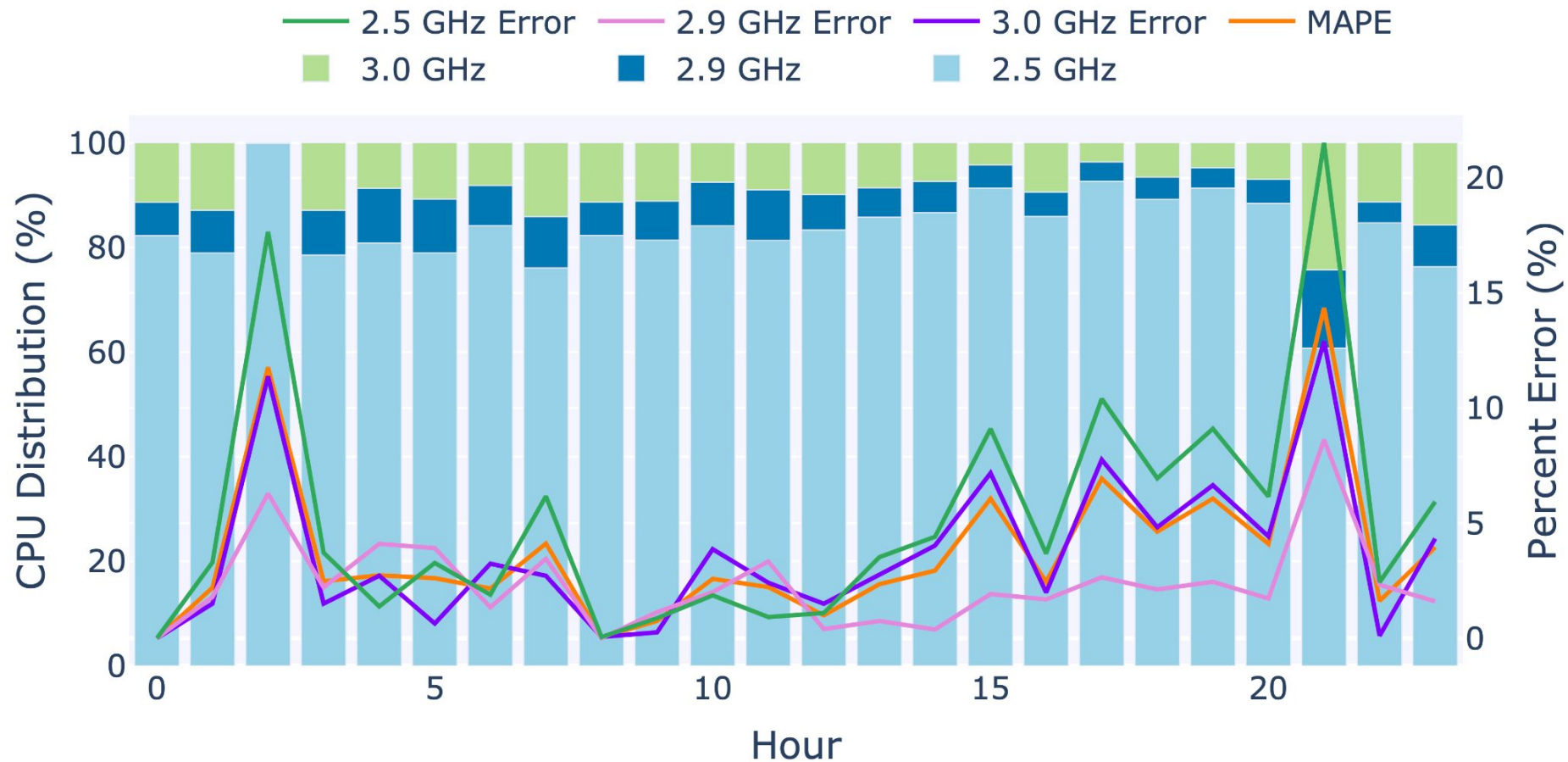








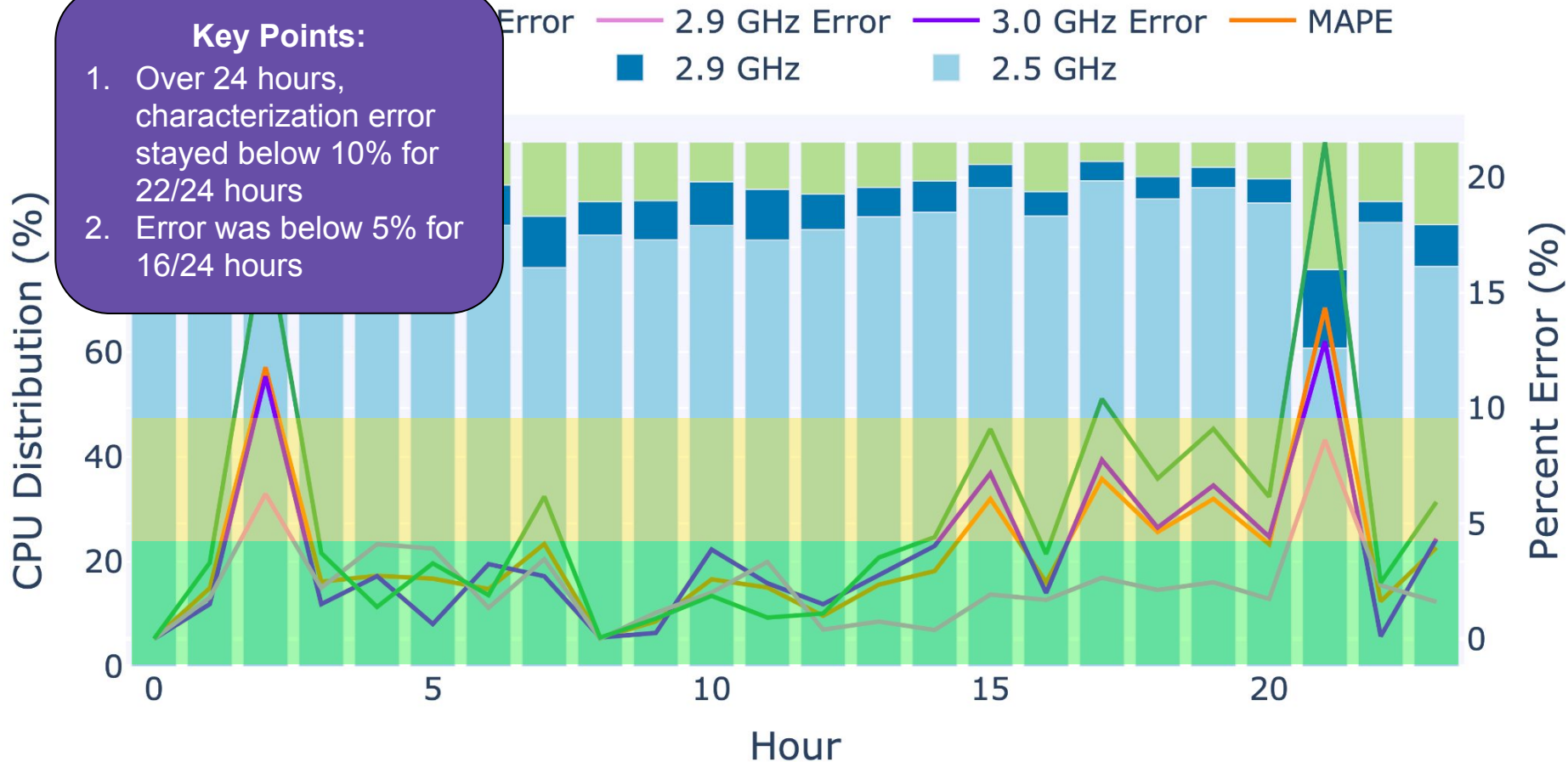
**RQ-1 (Part 2): How does
characterization accuracy
change over time?**



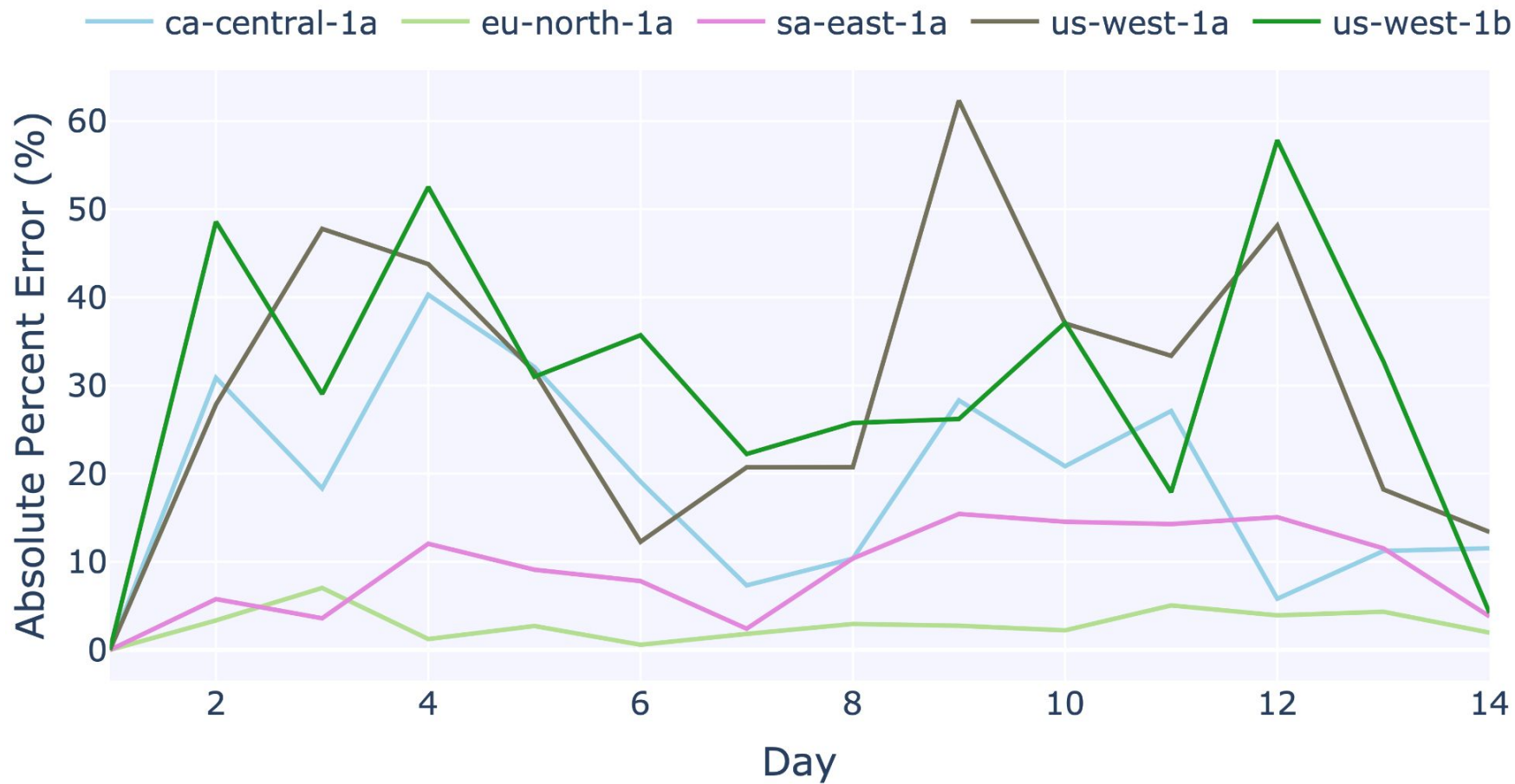
- Characterization accuracy over 24 hours on us-west-1b

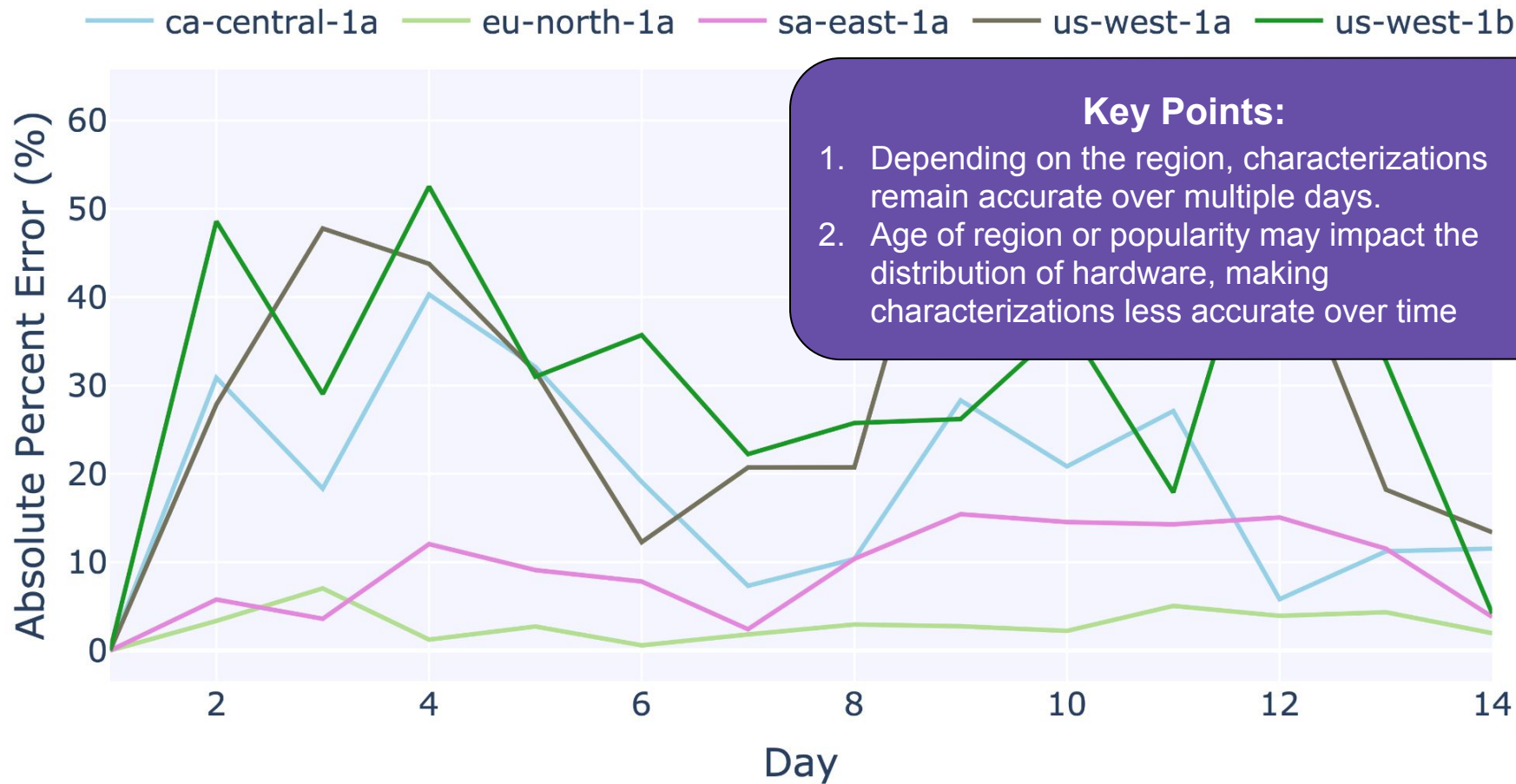
Key Points:

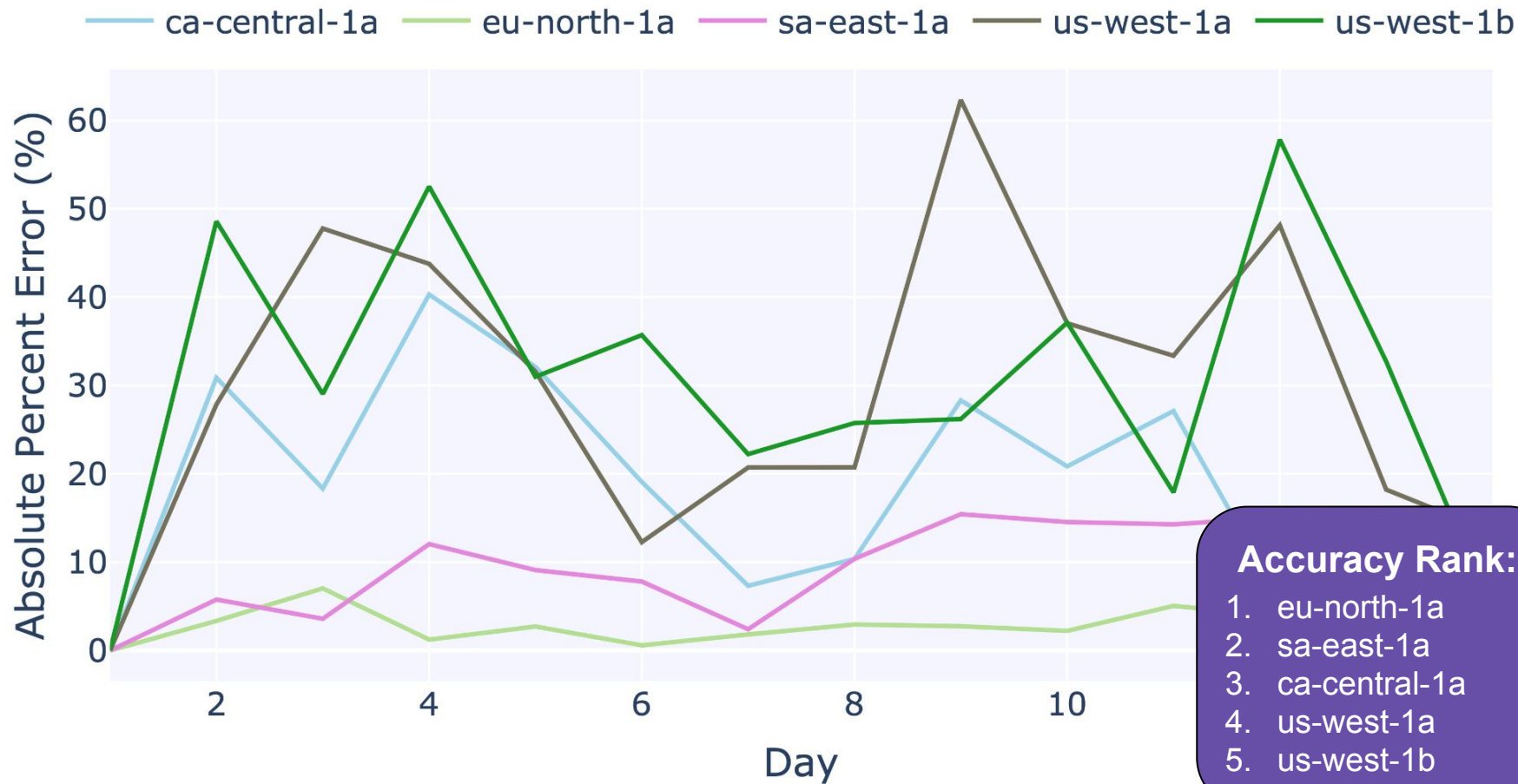
1. Over 24 hours, characterization error stayed below 10% for 22/24 hours
2. Error was below 5% for 16/24 hours



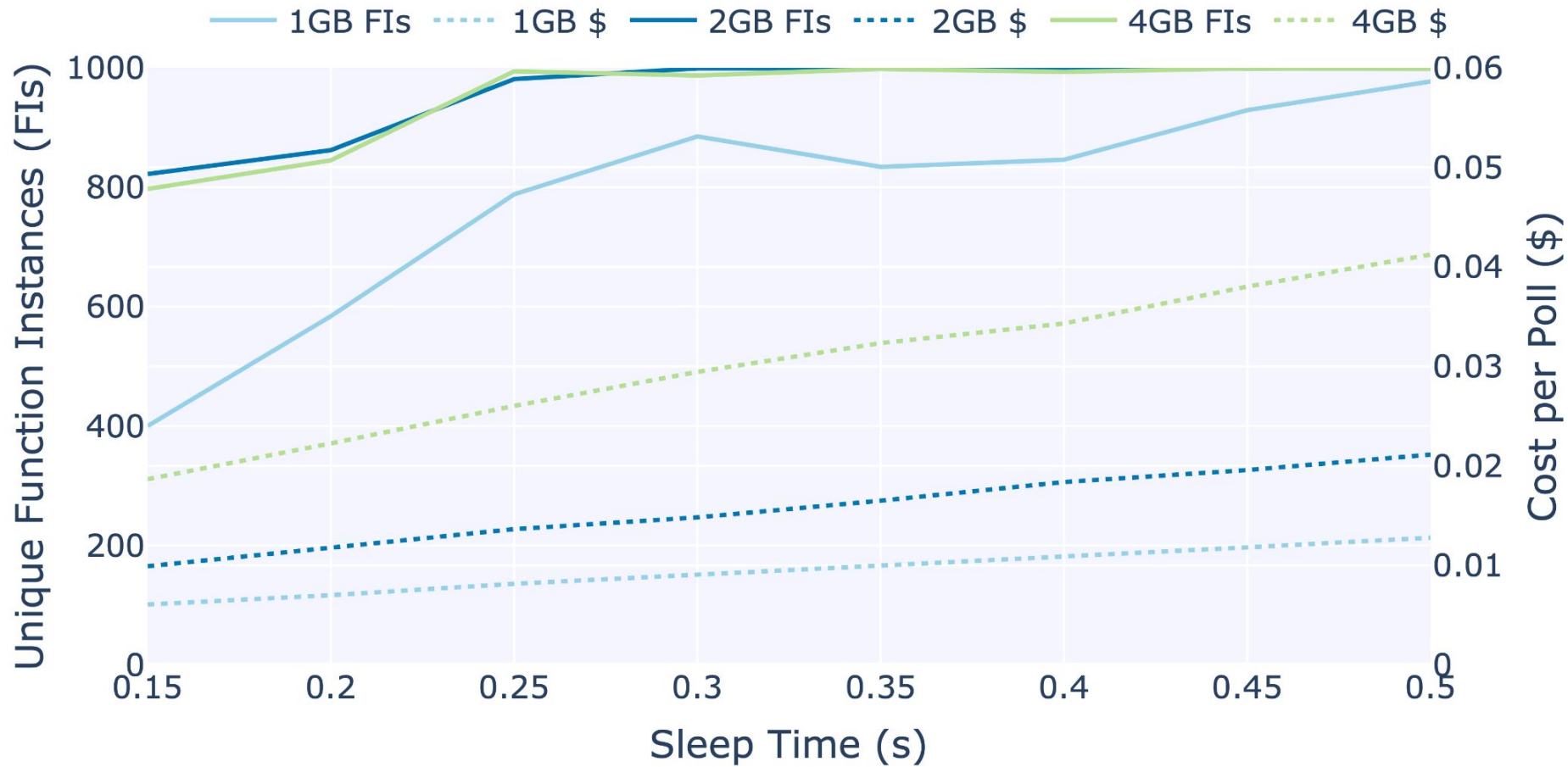
- Characterization accuracy over 24 hours on us-west-1b

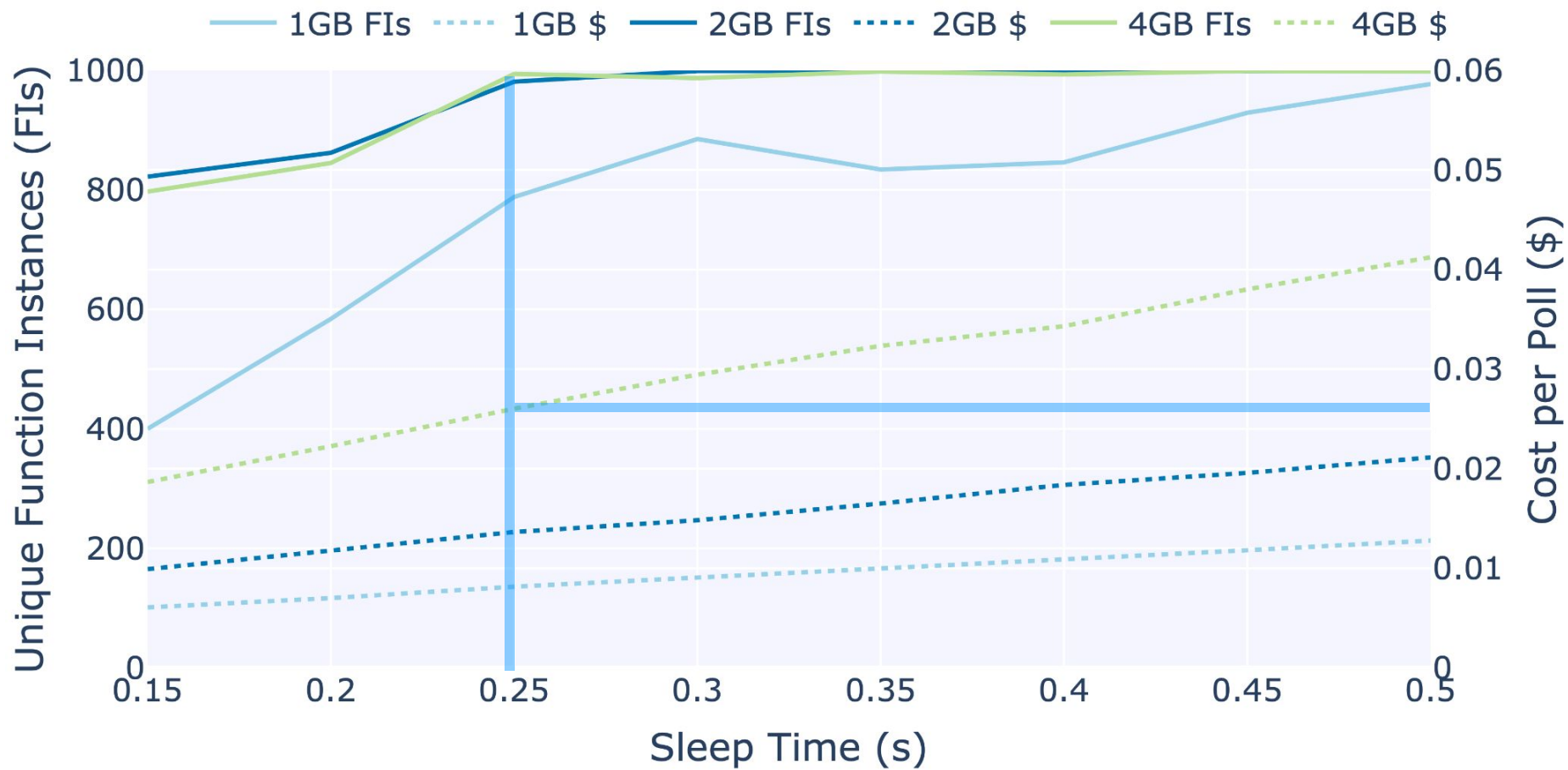


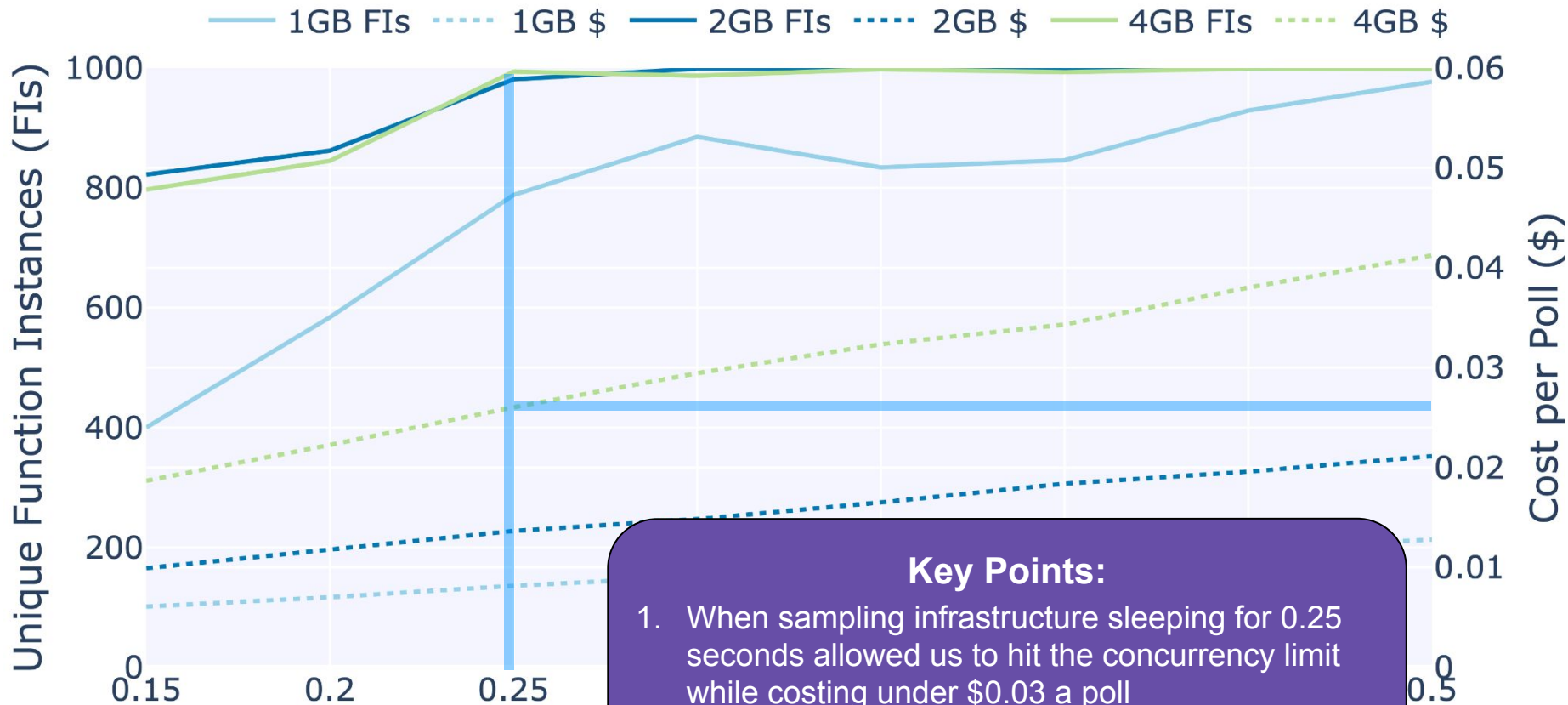




RQ-2: How many samples are required to accurately infer the hardware pool? How can we balance the cost versus accuracy?

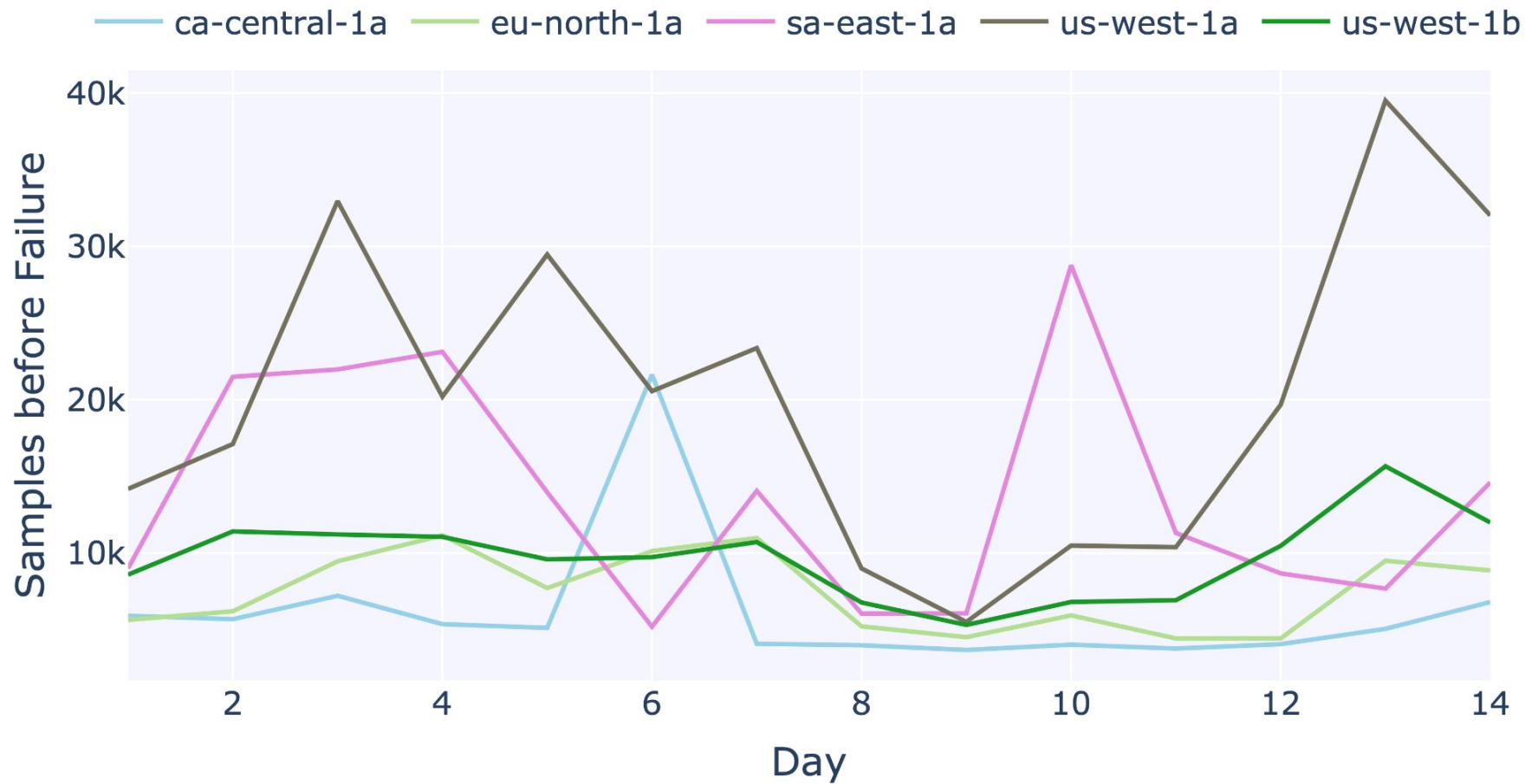


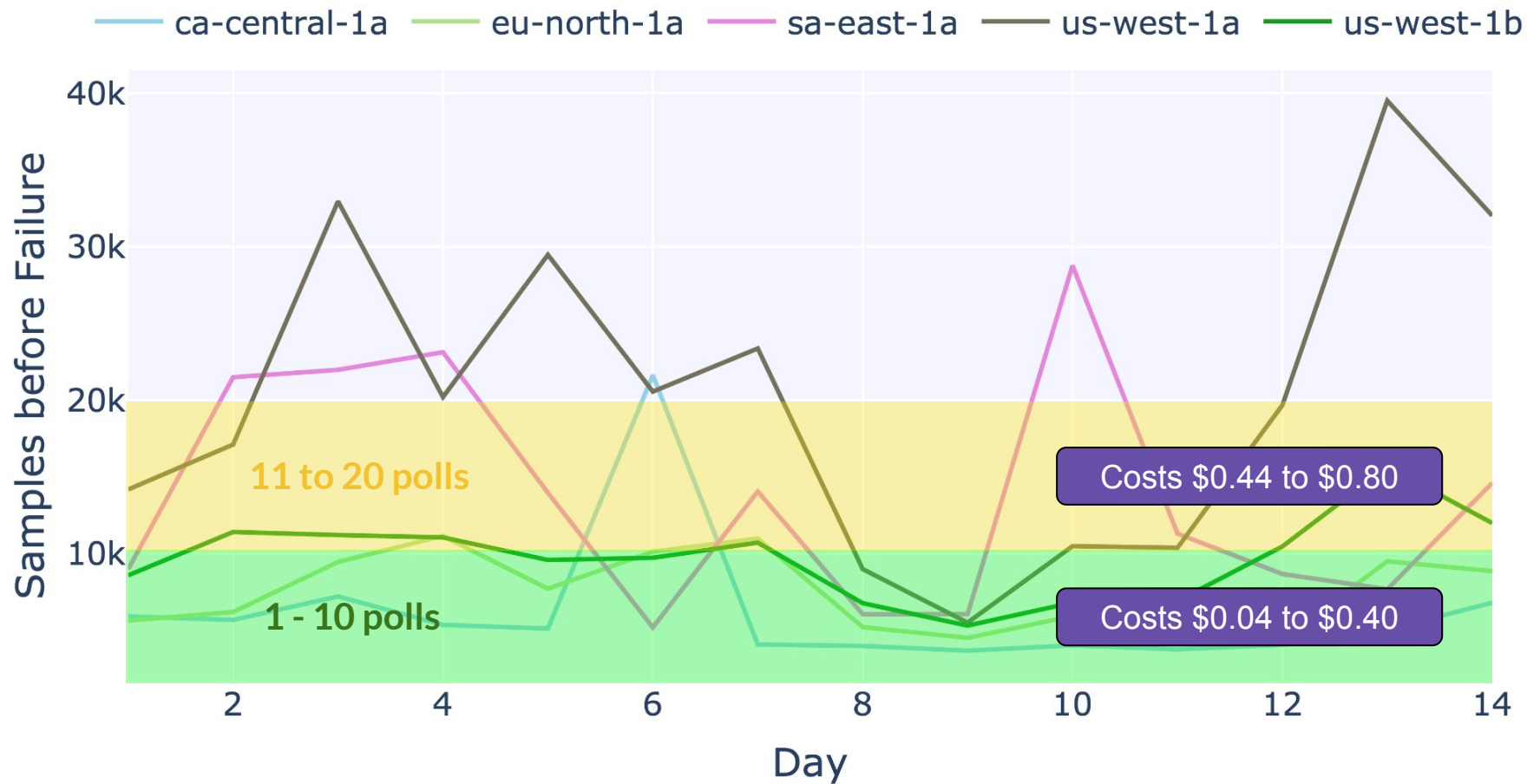


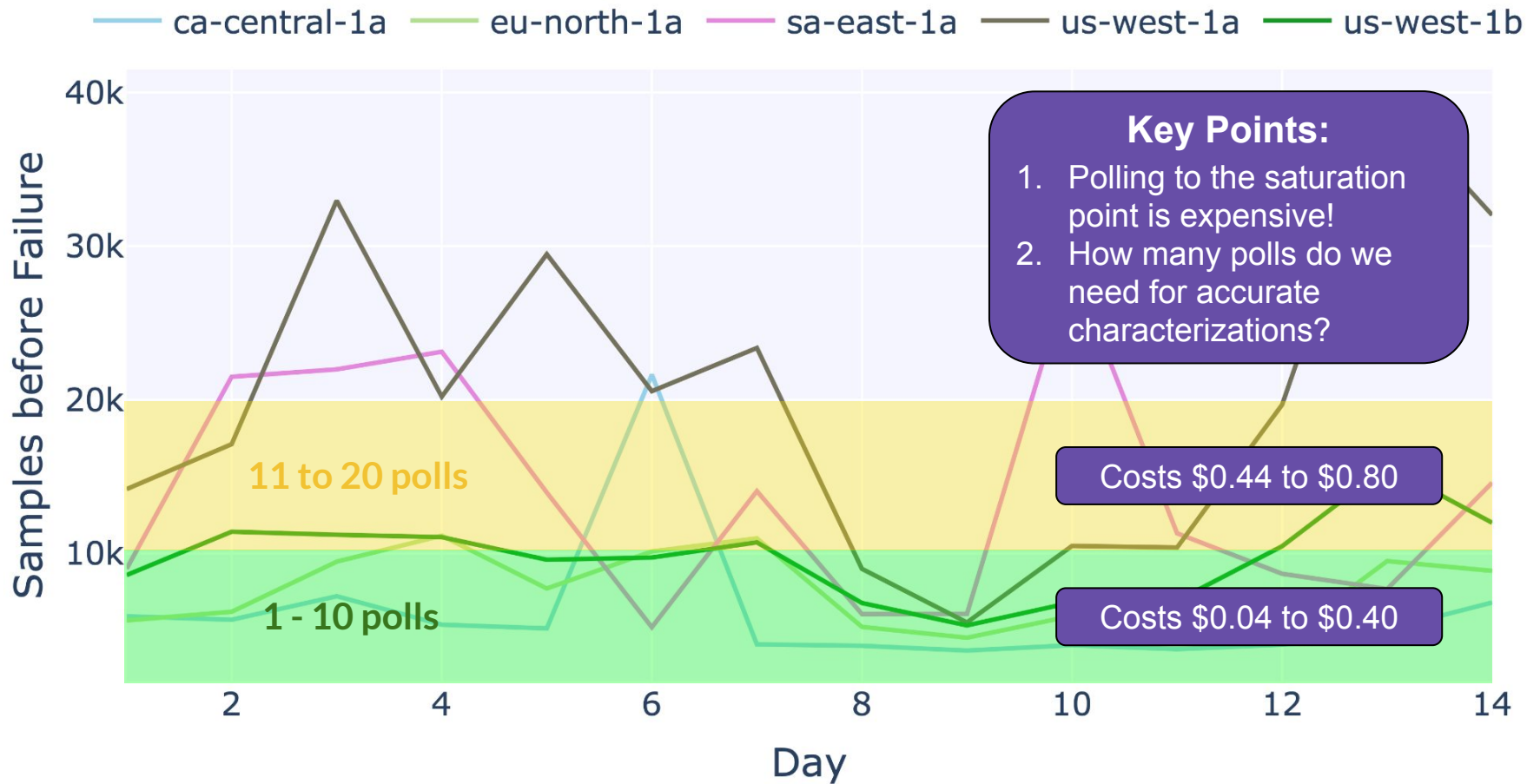


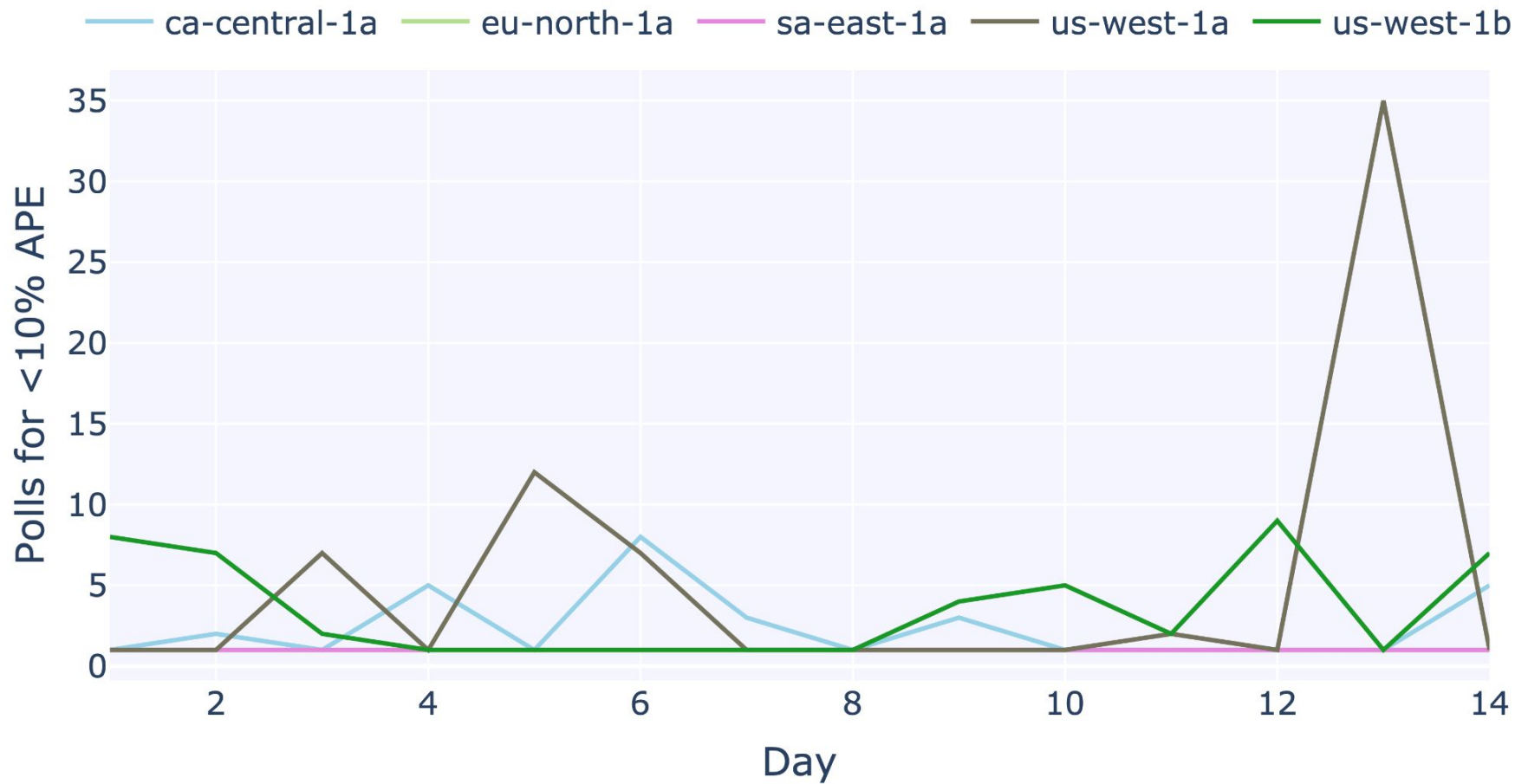
Key Points:

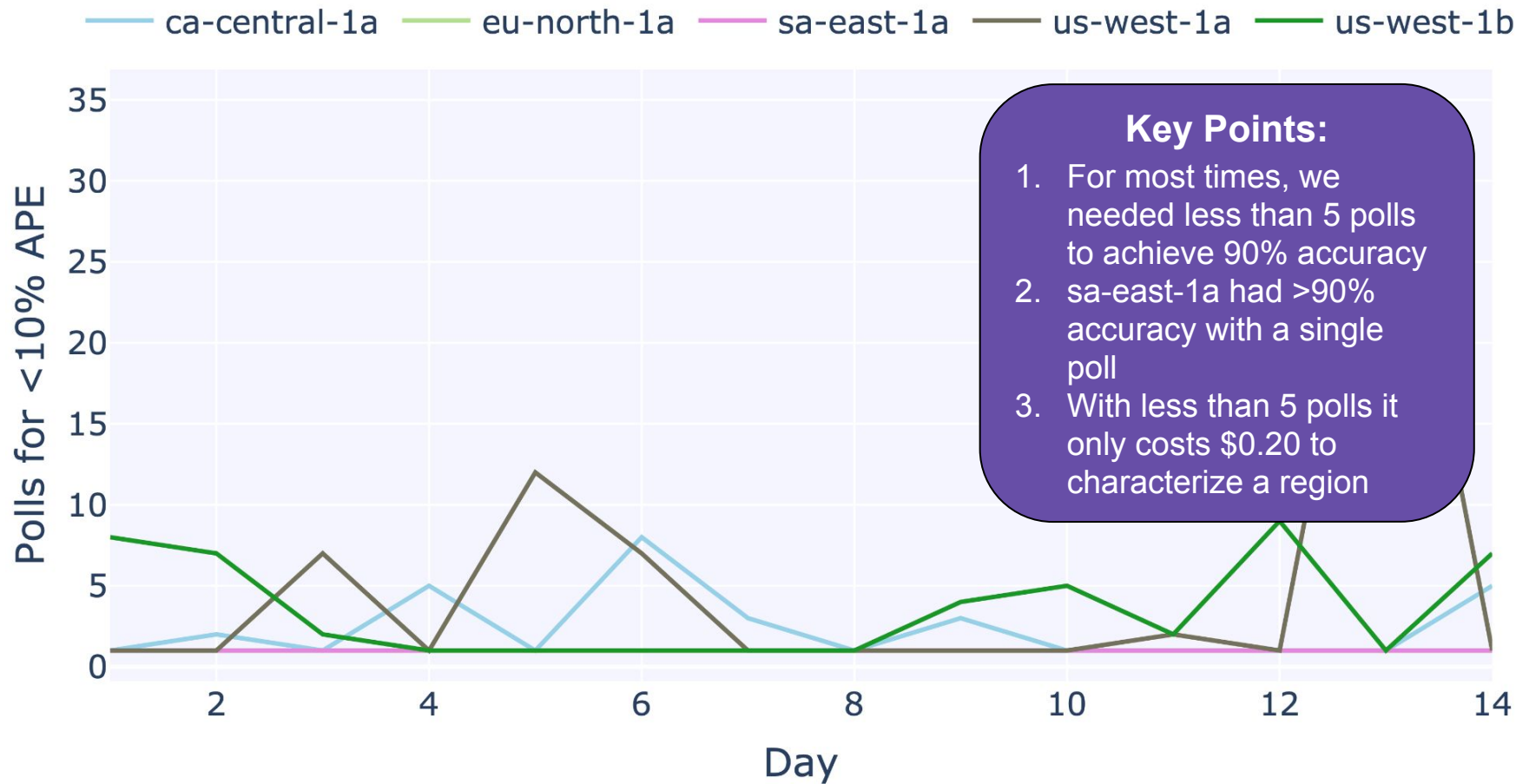
1. When sampling infrastructure sleeping for 0.25 seconds allowed us to hit the concurrency limit while costing under \$0.03 a poll
2. For later experiments we used 10 GBs memory settings, resulting in ~\$0.04 cost per poll



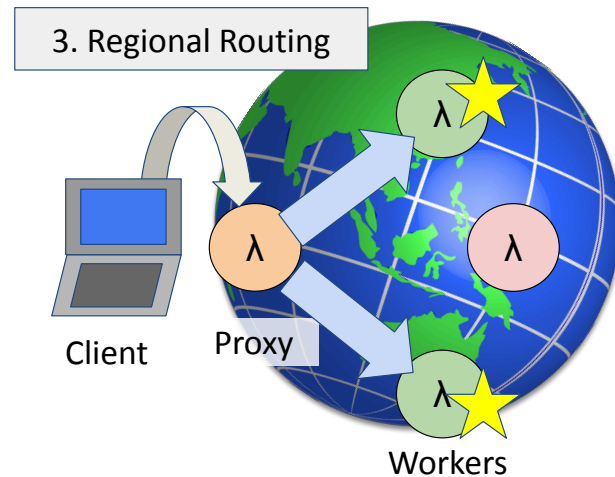
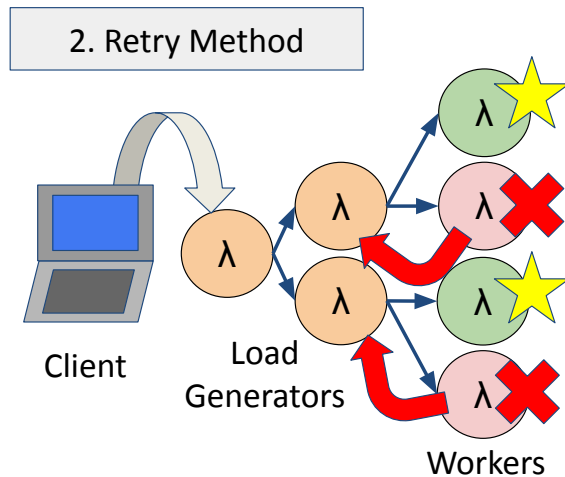
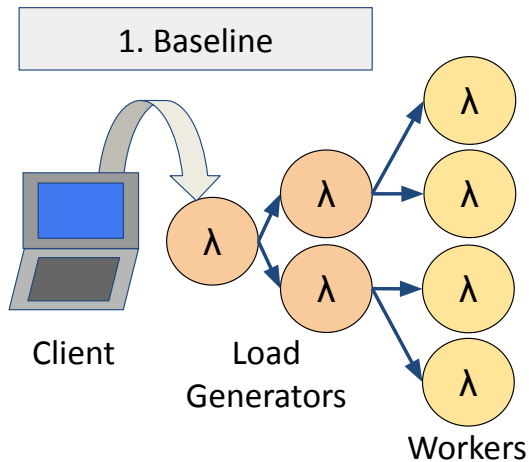








**RQ-3: To what extent are
runtime and cost
improvements possible by
targeting specific instructure?**

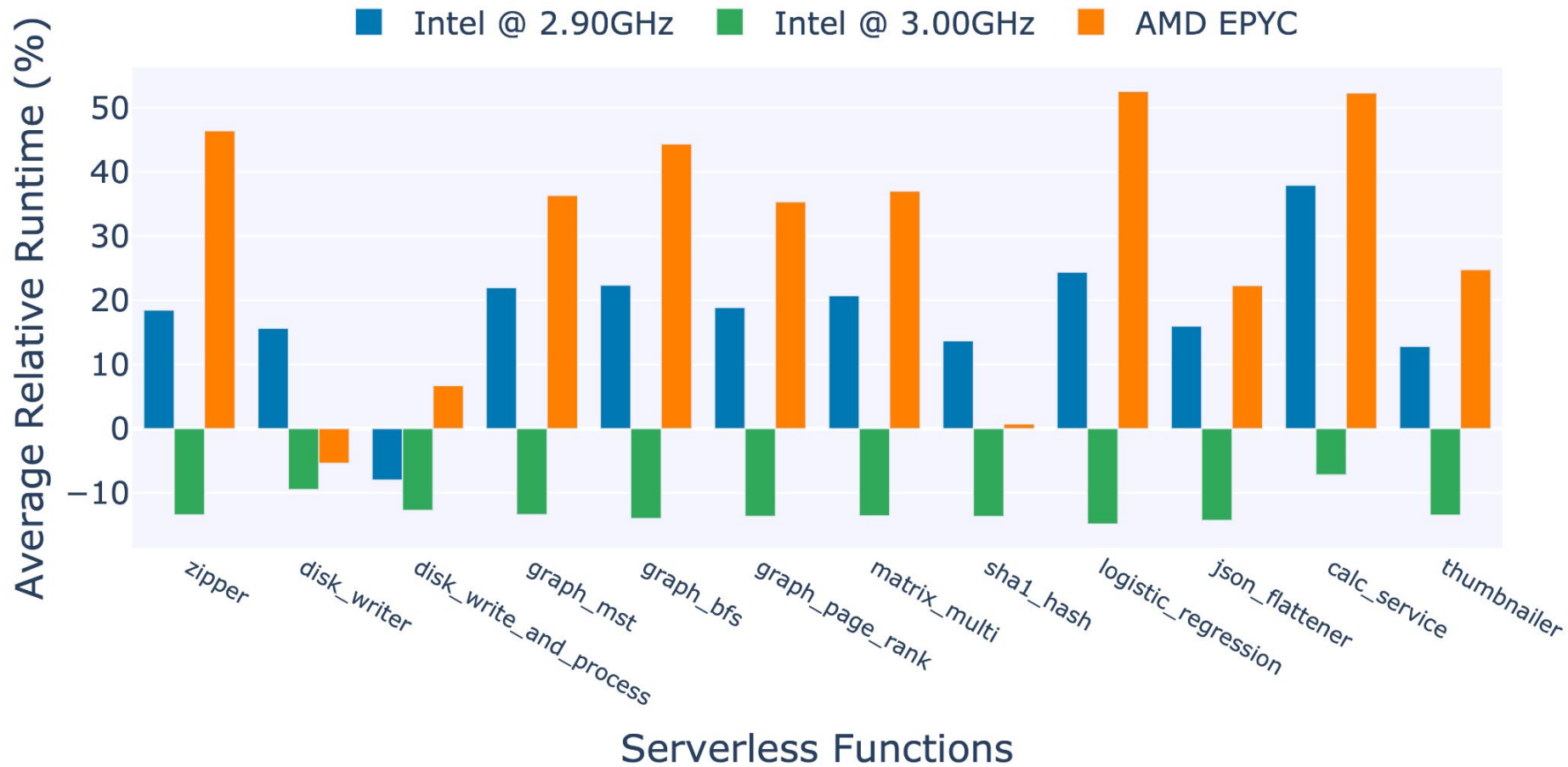


Retry Logic

1. **Retry Slow:** Invocations retry on the two slowest CPUs
2. **Focus Fastest:** Invocations retry on everything but the fastest CPU

Hybrid Approach

Combine both Retries and Regional Routing



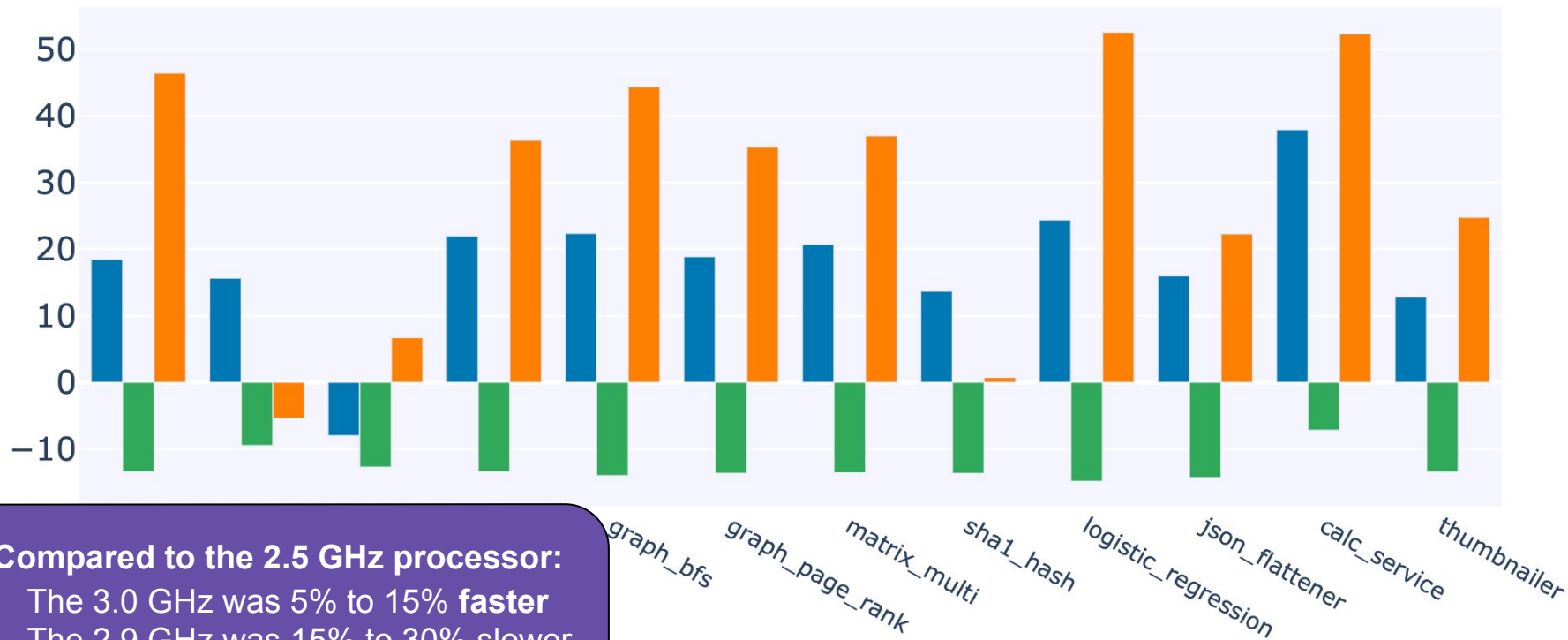
- Runtime relative to the most common Intel 2.5 GHz processor observed on AWS Lambda



- Runtime relative to the most common Intel 2.5 GHz processor observed on AWS Lambda

Average Relative Runtime (%)

■ Intel @ 2.90GHz ■ Intel @ 3.00GHz ■ AMD EPYC



Compared to the 2.5 GHz processor:

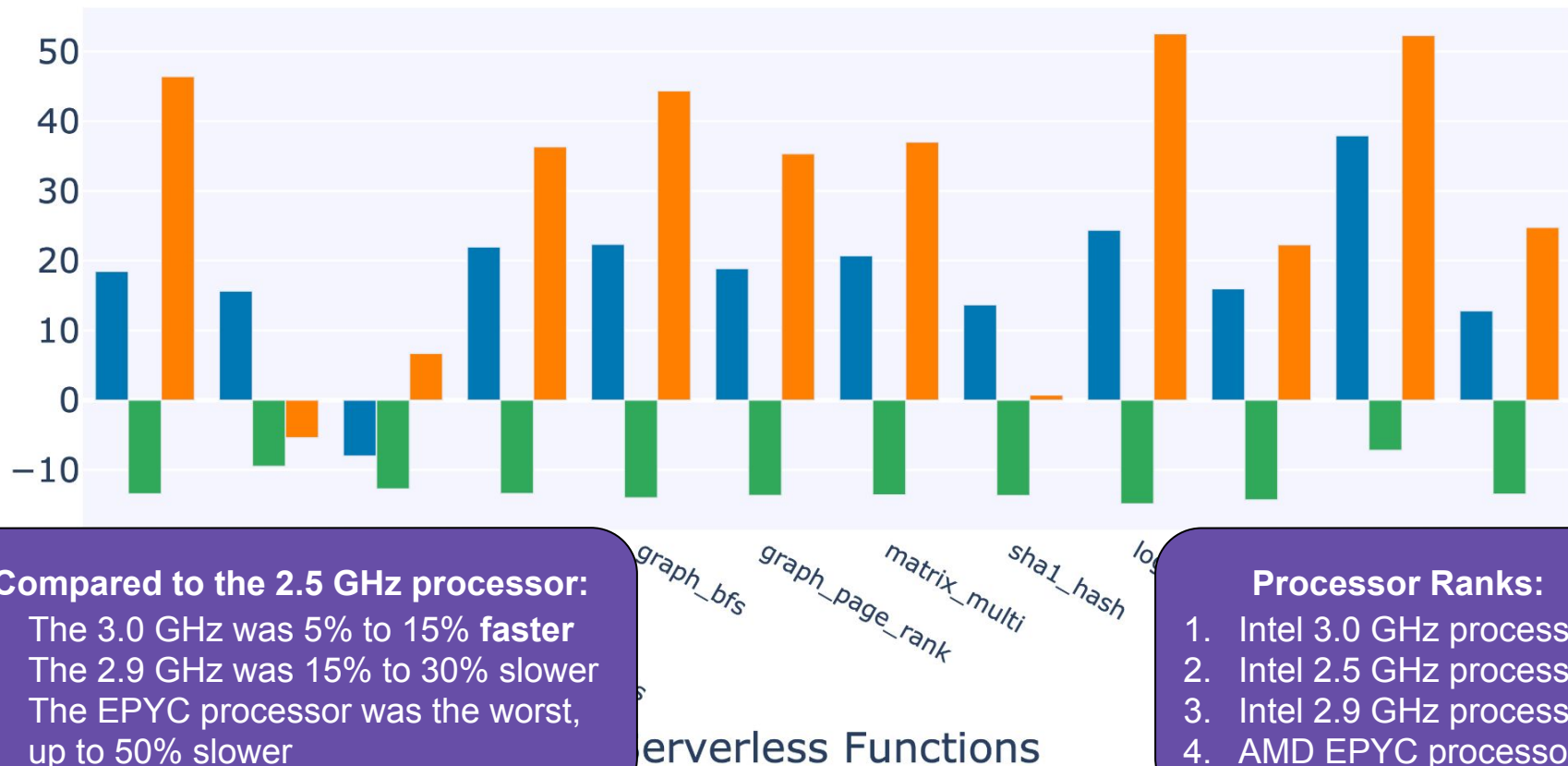
1. The 3.0 GHz was 5% to 15% **faster**
2. The 2.9 GHz was 15% to 30% **slower**
3. The EPYC processor was the worst, up to 50% **slower**

Serverless Functions

- Runtime relative to the most common Intel 2.5 GHz processor observed on AWS Lambda

Average Relative Runtime (%)

■ Intel @ 2.90GHz ■ Intel @ 3.00GHz ■ AMD EPYC



Compared to the 2.5 GHz processor:

1. The 3.0 GHz was 5% to 15% **faster**
2. The 2.9 GHz was 15% to 30% slower
3. The EPYC processor was the worst, up to 50% slower

Processor Ranks:

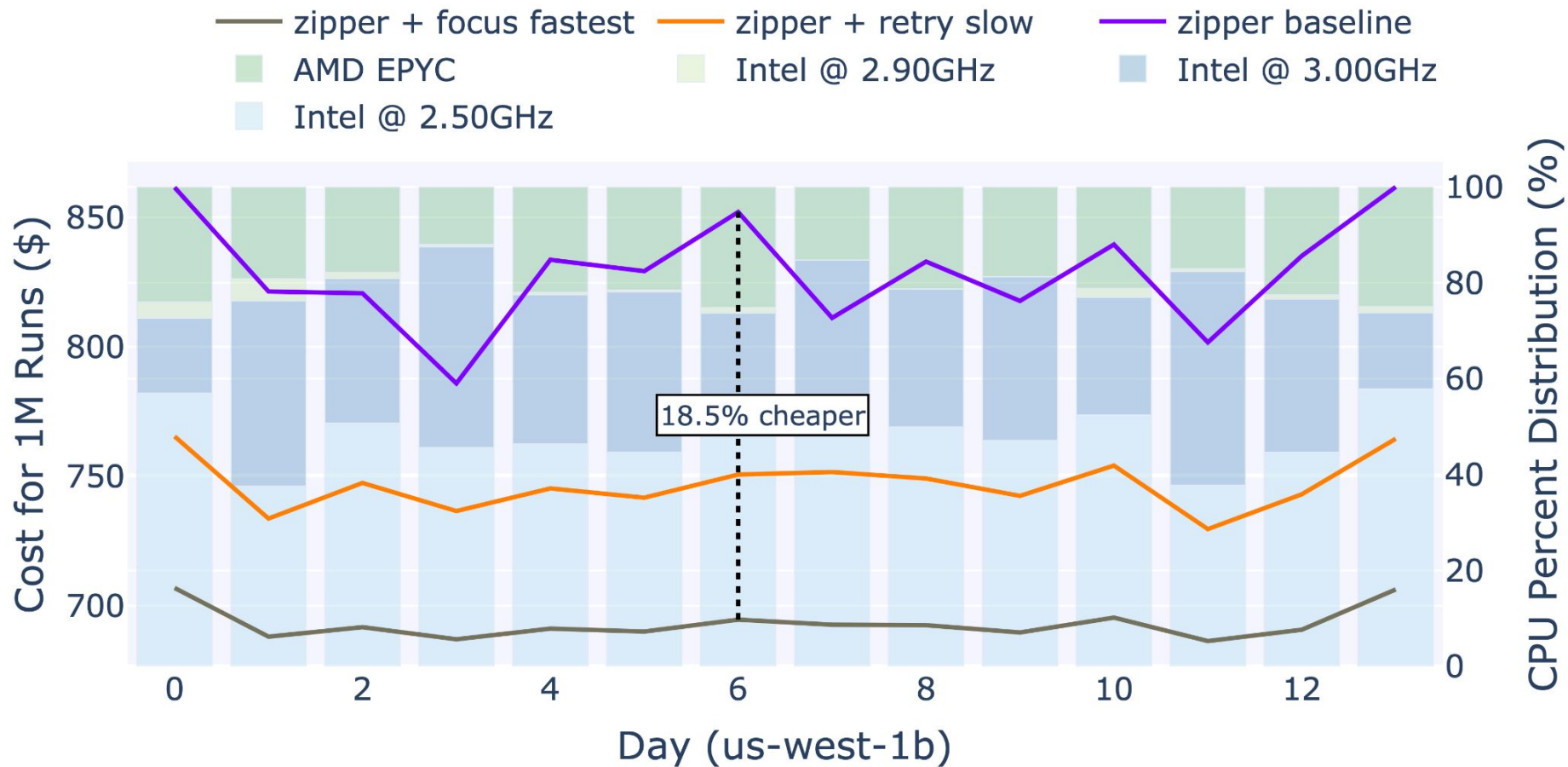
1. Intel 3.0 GHz processor
2. Intel 2.5 GHz processor
3. Intel 2.9 GHz processor
4. AMD EPYC processor

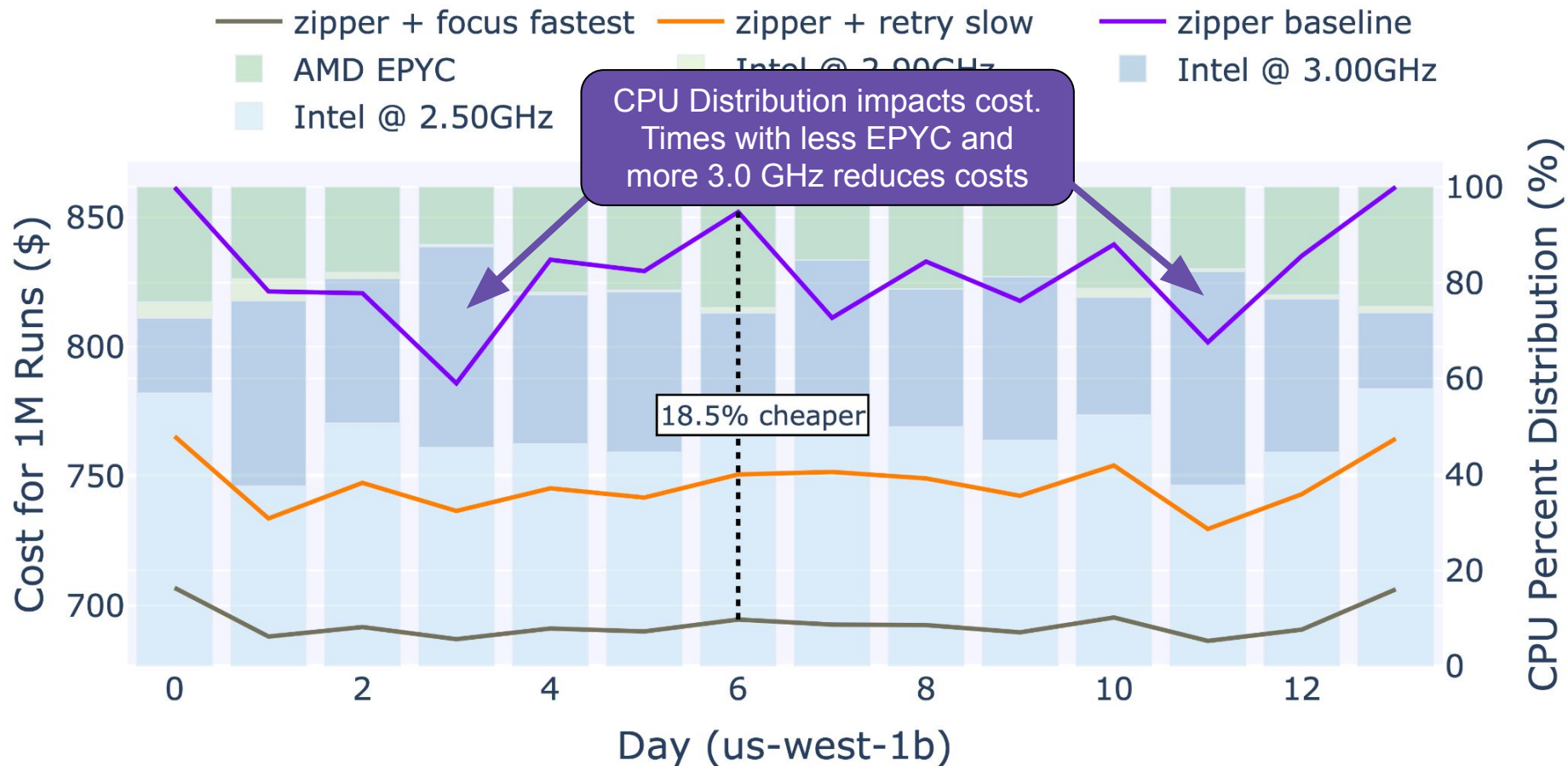
serverless Functions

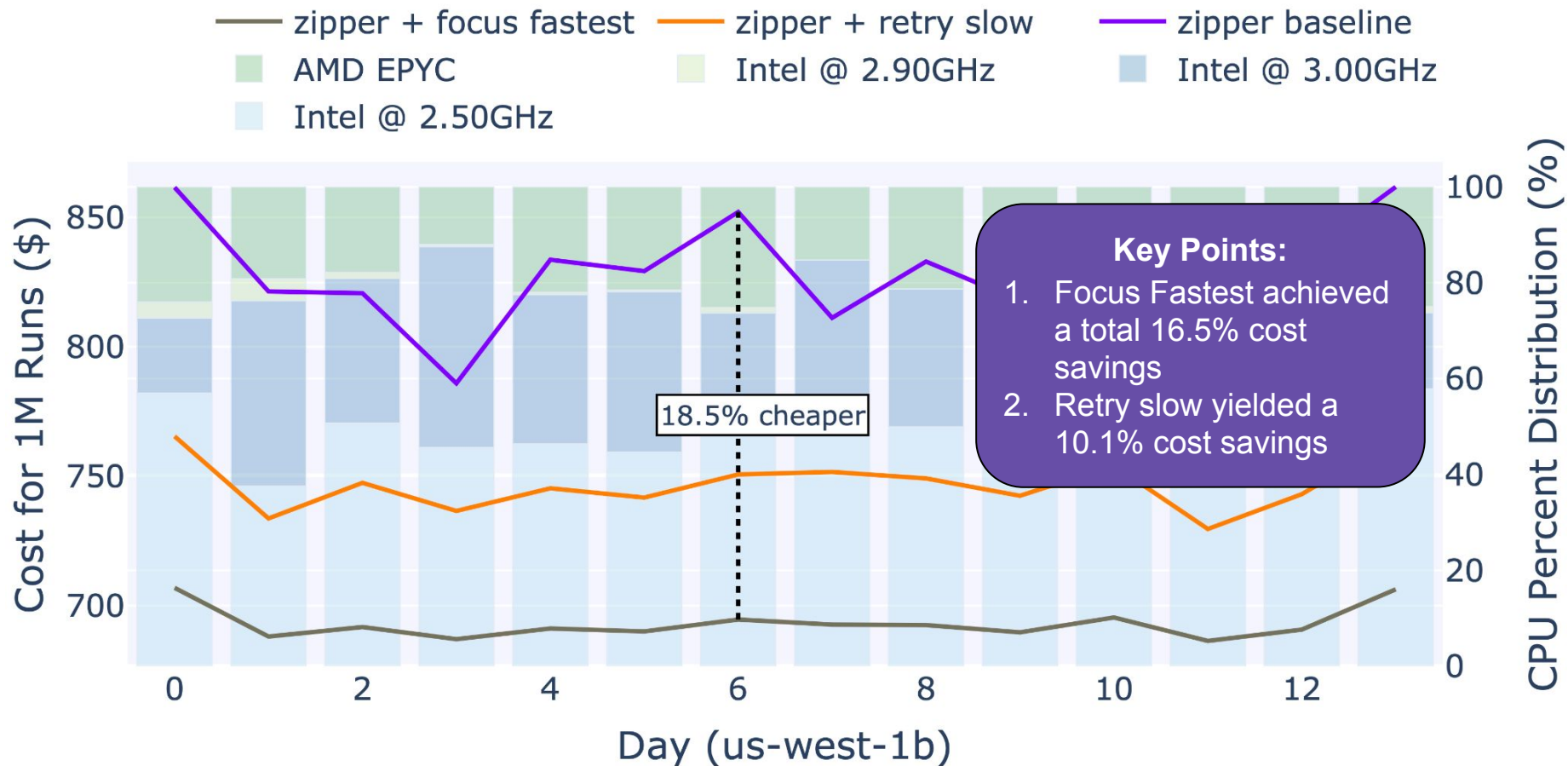
- Runtime relative to the most common Intel 2.5 GHz processor observed on AWS Lambda



Retry Methods

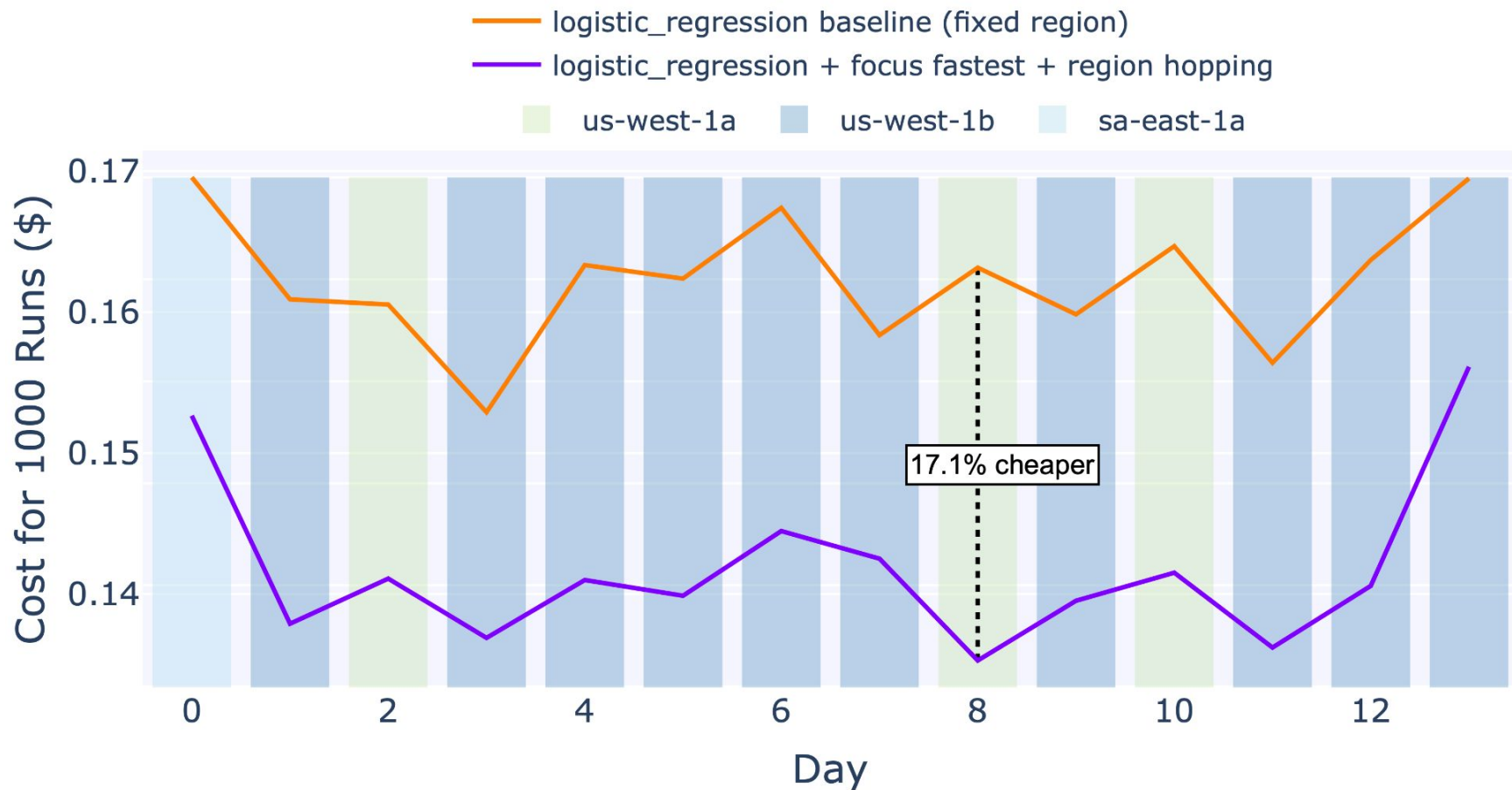




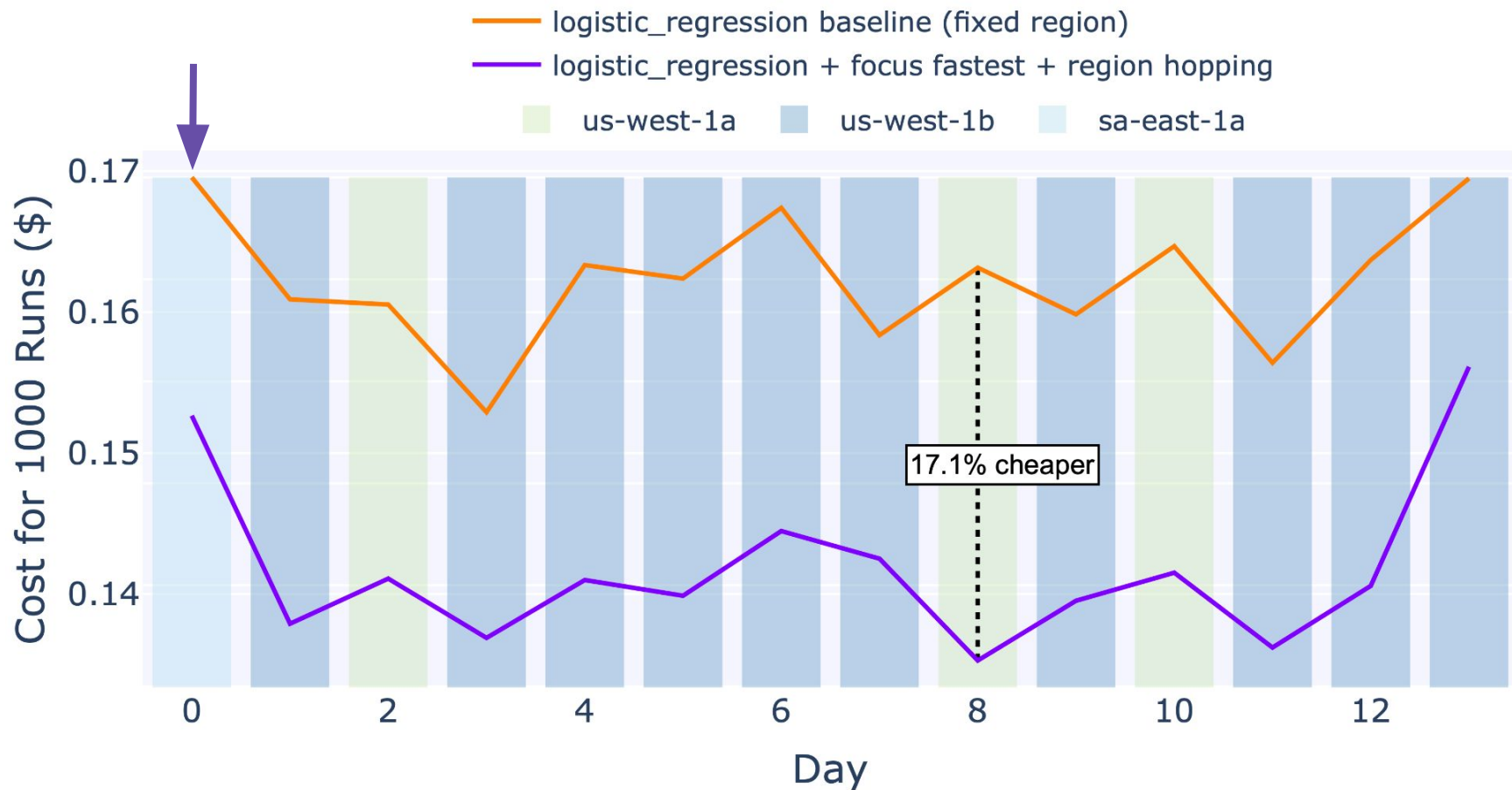




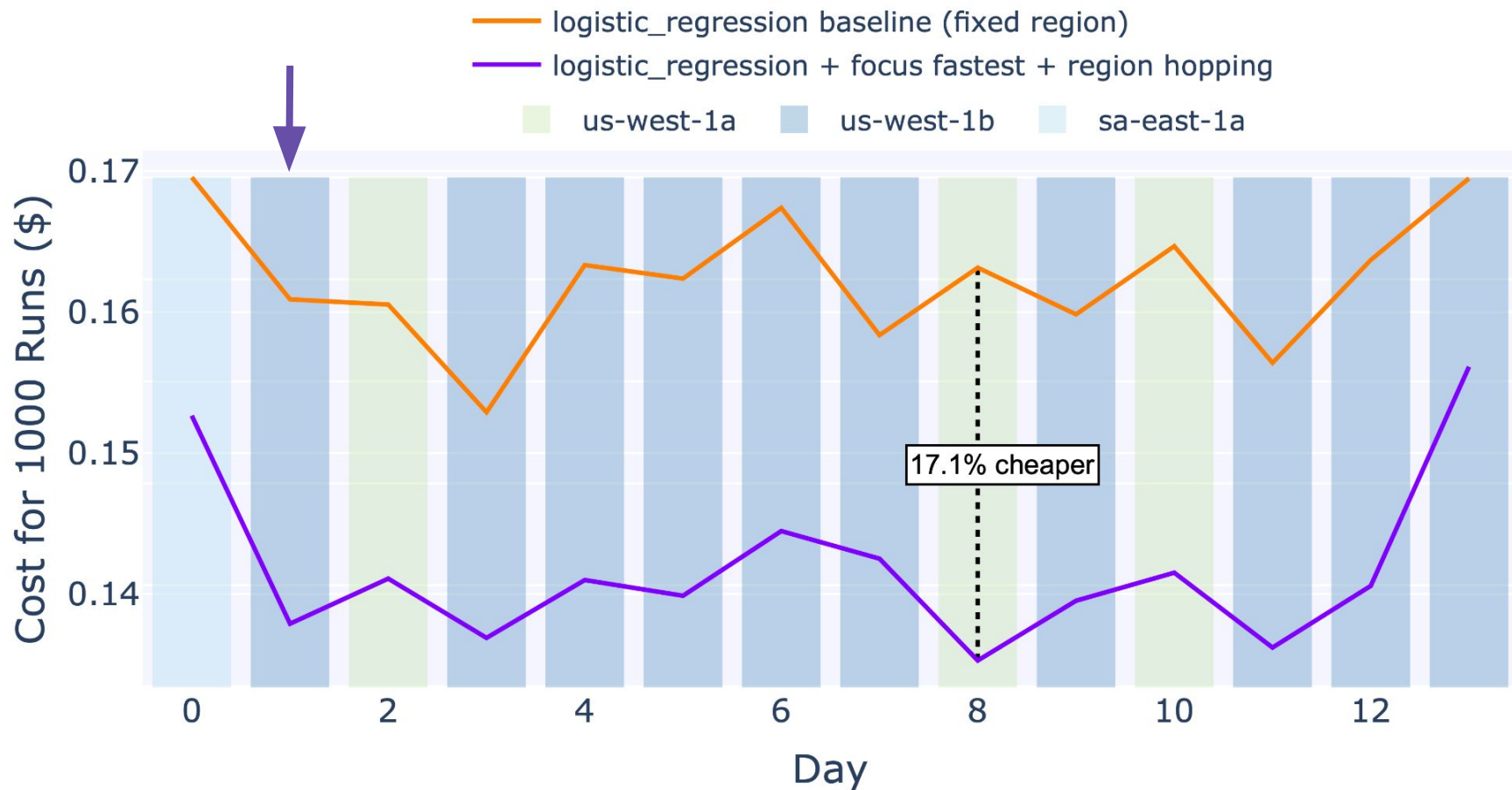
Region Hopping



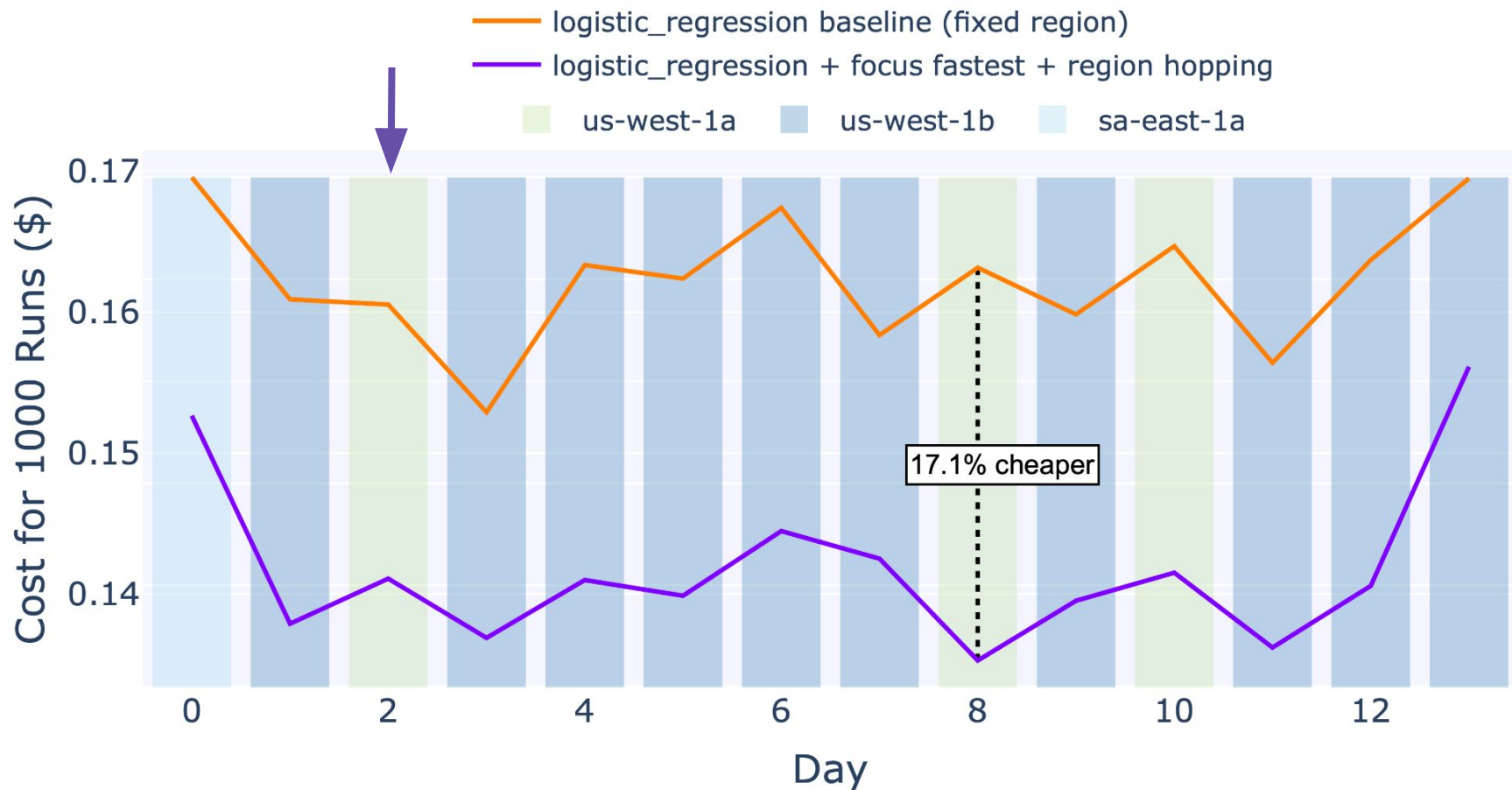
- Baseline fixed region of us-west-1b



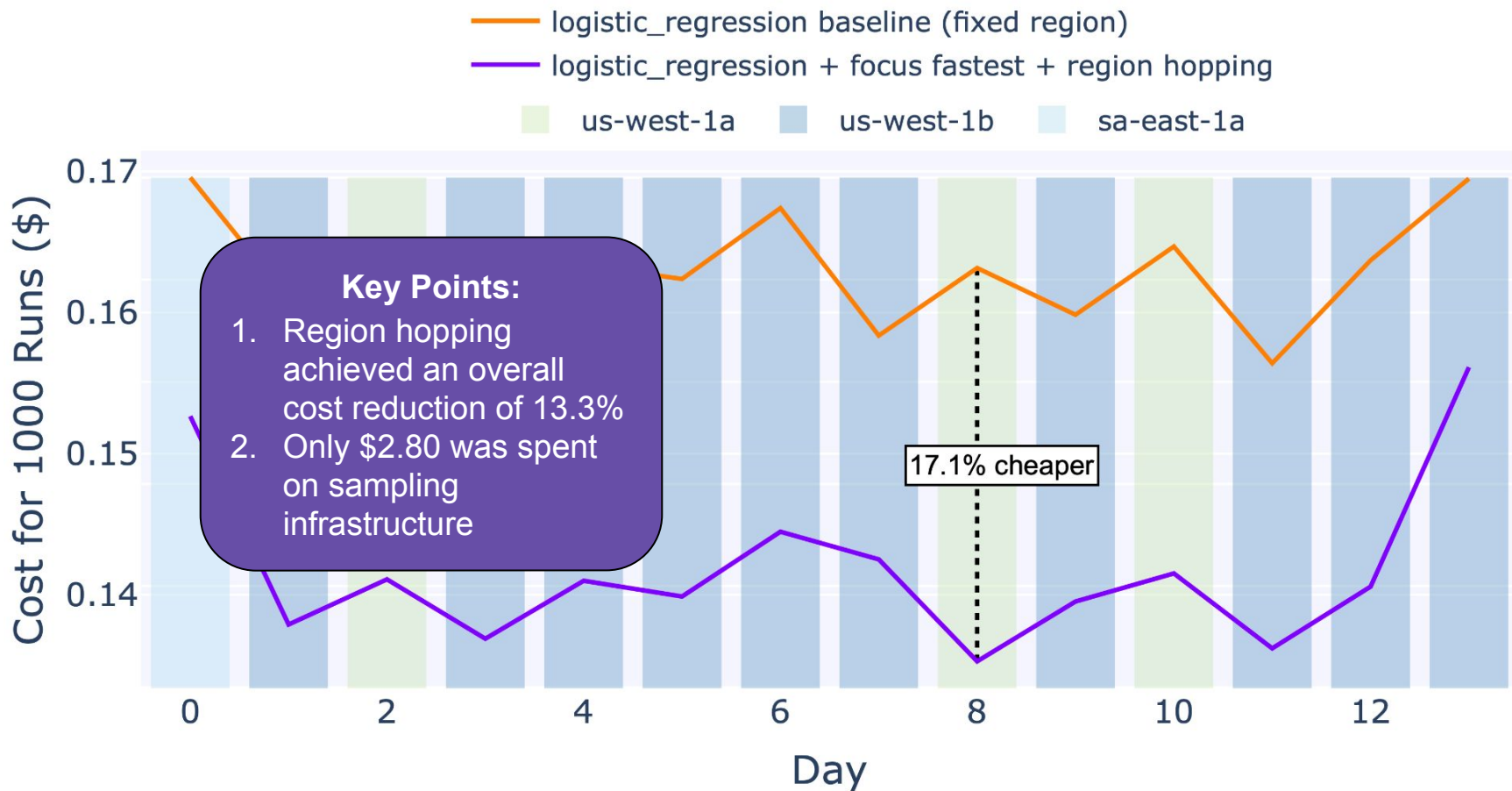
- Baseline fixed region of us-west-1b



- Baseline fixed region of us-west-1b



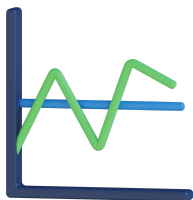
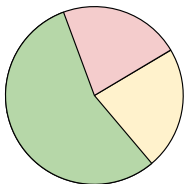
- Baseline fixed region of us-west-1b



- Baseline fixed region of us-west-1b



Conclusions



Key Results

- **RQ-1 (Infrastructure Variation):**
 - Diverse CPU characterizations were observed across 41 regions
 - Our sampling technique appears to be able to saturate entire availability zones creating accurate hardware characterizations
- **RQ-2 (Infrastructure Characterization):**
 - Characterized availability zones for only \$0.04 with 95% accuracy
- **RQ-3 (Performance Enhancement):**
 - **Retry Slow** resulted in a consistent 10.1% reduction in runtime and cost
 - **Focus Fastest** had up to 18.5% lower costs, averaging 16.5%
 - Compared to a baseline of running in us-west-1b, the **Region Hopping** hybrid approach averaged a cost savings of 10.0% up to 18.2%



Thank You!

Any Questions?

Presented by Robert Cordingly – rcording@uw.edu
All source code available at <https://github.com/wlloyduw/SAAF>

This research has been supported by the University of Washington Carwein-Andrews Distinguished Fellowship and AWS Educate Cloud Credits.