



netidee

STIPENDIEN

Container scheduling using AI-based workload
characterization for heterogeneous computer
clusters

Zwischenbericht | Call 15 | Stipendium ID 5237

Lizenz: CC-BY

Inhalt

1	Einleitung.....	3
2	Status.....	3
2.1	Meilenstein 1 - <Kurzbezeichnung>.....	3
2.2	Meilenstein 2 - <Kurzbezeichnung>.....	3
2.3	Meilenstein 3 - <Kurzbezeichnung>.....	4
2.4		4
3	Zusammenfassung Planaktualisierung.....	4

1 Einleitung

Meine Arbeit ist in den Bereichen Edge bzw. Serverless Computing angesiedelt und verschmilzt diese neuartigen Gebiete zu Serverless Edge Computing.

Bei Edge Computing handelt es sich um ein Paradigma, das entstanden ist um einerseits die riesigen Datenmengen zu bewältigen die vom Internet Of Things produziert wird (bspw. Im Kontext von Smart Cities) und andererseits um Anwendungen zu realisieren die eine sehr niedrige Latenz benötigen. Augmented/Virtual Reality und Autonomous Driving sind Anwendungsfälle, die in diese Kategorie fallen.

Serverless Computing hingegen siedelt sich im Bereich der Modelle für Deployments an. Durch Cloud Computing entstanden verschiedene Modelle um Applikationen an die jeweiligen Provider zu übertragen. Durch die populär gewordene Technik von Containern, wurde vor allem Platform-as-a-Service zunehmend beliebter. Container werden von Applikationsentwicklern erstellt und enthalten das Programm, sowie alle Abhängigkeiten. Sie sind daher ähnlich zu Virtual Machines, aber haben im Gegensatz einen weitaus geringeren Performance Overhead.

Speziell widme ich mich der Verbindung aus diesen neuartigen Gebieten und lege Fokus auf die Platzierung von Applikationen in solchen Infrastrukturen.

Weiters beschäftige ich mich mit neuartigen Hardwarebeschleunigern und versuche ein Matching zwischen Applikationsanforderungen und Gerätefähigkeiten zu bestimmen.

Dieses Matching soll auf Basis von Workload Characterization ermöglicht werden und ist insofern wichtig, weil es im Edge Bereich eine große Anzahl von verschiedenen Geräten gibt, die alle unterschiedlich "stark" sind und deshalb braucht der Scheduler zusätzliches Wissen um bestmögliche Platzierungen treffen zu können.

Ich befinde mich derzeit schon in einem sehr fortgeschrittenen Zustand, da ich mit meiner Arbeit schon im Frühling/Sommer begonnen habe.

Zu der Zeit war die größte Herausforderungen eine mögliche und anwendbare Lösungsstrategie zu entwickeln. Ich habe mich im Laufe der Zeit gegen einen RL-Agenten entschieden und verwende nun Machine Learning und habe ein Optimierungsproblem formuliert, das ich mit einem genetischen Algorithmus löse.

Weiters wurde das Anwendungsgebiet festgelegt: Plazieren von Anwendungen (Containern) mit Hilfe einer Function-as-a-Service Plattform mit Fokus auf AI Pipelines (Trainieren, Preprocessing und Inferenz)

Ich bin an einem Punkt angelangt, an dem alle notwendigen Vorkehrungen getroffen wurden, um die finale Evaluierung auszuführen.

2 Status

2.1 Meilenstein 1 – Profiling/Fertigstellung der Simulation/Beginn Zusammenschrift

Kurzbeschreibung der Haupttätigkeiten

Dieser Meilenstein umfasste die notwendigen Schritte um eine objektive Evaluierung zu ermöglichen.

Einerseits wurden für die gewählten Testapplikation Profiling Daten gesammelt, die der Simulator verwendet (mehr dazu hier: [1]).

Weitere Tätigkeiten zur Fertigstellung der Simulation umfassten:

- Implementierung einer plausiblen Topologie (angelehnt an eine Urban Sensing Szenario)
- Erstellung einer plausiblen Requestverteilung (= Simulation der Users)
- Fertigimplementierung einer Function-as-a-Service Plattform (mit Hilfe der Simulation, keine echte Plattform, siehe [2]), die zwar auf OpenFaaS basiert – Hardwarebeschleuniger aber als First-Class Citizens pflegt. OpenFaaS erlaubt es Usern nicht in einer Funktion verschiedene Arten dieser zu gruppieren. Deshalb musste ich ein System entwickeln, welches Usern die Möglichkeit gibt verschiedene Funktionen zu einer einzigen zu gruppieren.

Beispielsweise Objekterkennung, die wahlweise mit CPU oder GPU ausgeführt wird

Weiters habe ich mit der Niederschrift begonnen. Besonders wichtig war/ist mir bisher einerseits einen guten Hintergrund zu formulieren, sowie die verschiedenen Lösungen zu präsentieren um auf etwaige „Denkfehler“ vor der finalen Evaluierung zu stoßen.

Erkenntnisse zur Vorgangsweise

Da ich schon länger mit dieser Arbeit beschäftigt bin, hat sich die Vorgangsweise als ausreichend gut durchdacht entpuppt. Dies waren vorerst die letzten Schritte um meine Evaluierung auszuführen und sollte ich nicht noch auf Probleme stoßen, waren es die letzten die mich vor meinen Ergebnissen trennen.

Kurzbeschreibung der erreichten Ergebnisse

Die Ergebnisse bestehen einerseits aus Rohdaten, Code und der Anfang meiner Diplomarbeit.

Die Rohdaten bestehen einerseits aus Traces sowie Telemetriedaten. Jede Testapplikation wurde mehrfach auf jedem Geräte (insgesamt 9) ausgeführt und verschiedene Zeitstempel (Ankunft/Ende des Requests) mitespeichert. Dadurch kann ich mir bspw. Die Function

Execution Time (FET) berechnen, die meine Performancemetrik darstellt.

Telemetriedaten beinhalten CPU Usage, I/O Bytes, Block I/O Bytes, RAM und GPU Usage.

Anhand der Telemetriedaten wird ein Machine Learning Modell trainiert um mögliche Verschlechterung der Performance von Applikation aufgrund anderer auf denselben Geräten laufenden Apps zustande kommt.

Weiters werden sie in meinen Schedulererweiterungen verwendet um Resourcestatus zu verhindern.

Die Simulation konnte erfolgreich erweitert werden und eine prototypische FaaS Plattform mit Hilfe der *faas-sim* Komponenten implementiert werden.

In diesem System können Nutzer/innen nun verschiedenen Arten (GPU/CPU/TPU/Tflite) einer Funktion deployen und die Plattform entscheidet automatisch welche Art beim Skalieren verwendet wird.

Das Schreiben verläuft auch reibungslos. Der wichtige Hintergrund zu *Edge Intelligence* wurde fertiggestellt sowie die ersten Schritte meiner Herangehensweise dargelegt.

Besondere Erfolge/ Probleme

Derzeit läuft alles nach Plan. Probleme gab es, wie das leider bei solchen Geräten der Fall ist mit wenig Ressourcen, beim Profiling. Beispielsweise hat ein sehr schwacher Raspberry Pi3 manchmal Probleme mit diversen AI-Applikationen (Tensorflow).

Gab es große Abweichungen zum Plan? Warum?

Es gab keine Abweichungen.

[1]: <https://netidee.at/container-scheduling-using-ai-based-workload-characterization-heterogeneous-computer-clusters/about>

[2]: <https://github.com/edgerun/faas-sim/blob/performance-degradation/sim/faas.py#L154>

2.2 Meilenstein 2 – Erstellung ZB/Blogeintrag

Kurzbeschreibung der Haupttätigkeiten

Dieser Meilenstein widmet sich mit der Erstellung meines ersten Blogbeitrags, sowie den Zwischenbericht.

Kurzbeschreibung der erreichten Ergebnisse

Ich hab mit meinem Blogbeitrag versucht die Thematik meiner Arbeit (Edge Computing) für ein breites Publikum verständlich zu gestalten und habe Einblick in meine Evaluation gegeben, da das aufgrund fehlender Expertise und Standardbenchmarks ein weiteres Problem darstellt bei der Weiterentwicklung von Anwendungen/Plattformen in diesem Gebiet.

Den Blogbeitrag findet man hier: <https://netidee.at/container-scheduling-using-ai-based-workload-characterization-heterogeneous-computer-clusters/about>