

Cloud-Basierte Optimierung von Fahrzeugbetriebsstrategien durch Clustering mit Genetischen Algorithmen

Natalia Ogulenko¹, Sören Frey², Jens Nahm³ und Manfred Rössle⁴

Abstract: Die zunehmende Onlinevernetzung von Fahrzeugen eröffnet Automobilherstellern neue Wege, ihren rechtlichen Verpflichtungen zur Senkung der CO₂-Emissionen nachkommen und entsprechende Wettbewerbsvorteile realisieren zu können. Die für viele Fahrten übermittelten Telemetrie- und Streckendaten können analysiert und zur Senkung des Flottenverbrauchs verwendet werden. Wir stellen einen Ansatz vor, der diese Daten zur Optimierung von Fahrzeugbetriebsstrategien verwendet. In einer Cloud-Anwendung werden die übermittelten Daten anonymisiert gespeichert. Um die für bestimmte Charakteristiken, wie z.B. Motorleistung und Wetterbedingungen, am besten geeignete Betriebsstrategie aus den großen Datenmengen effizient bestimmen zu können, schränkt unser Ansatz den Suchraum mit einer Clusteranalyse ein. Wir verwenden hierzu ein Verfahren auf Basis von genetischen Algorithmen. Ein Cluster enthält eine Menge von möglichst ähnlichen Telemetrie- und Streckendaten. Die für eine bevorstehende Fahrt am besten geeignete Betriebsstrategie wird dann aus Daten des ähnlichsten Clusters bestimmt. Zukünftig könnten die so gefundenen optimalen Betriebsstrategien dynamisch vor Fahrtantritt zum jeweiligen Fahrzeug übertragen werden.

Keywords: Fahrzeugbetriebsstrategien, Optimierung, Clusteranalyse, Genetische Algorithmen

1 Einleitung

Durch die Onlinevernetzung von Fahrzeugen und die Verwendung von übermittelten Daten in Cloud-basierten Big Data Analysen ergeben sich vielfältige Potenziale in weiten Teilen der Automobilindustrie [KW14]. Dieses Papier beschreibt unseren Ansatz [Og16] zur Bestimmung einer optimalen Fahrzeugbetriebsstrategie pro bevorstehender Fahrt aus historischen Fahrten durch Clustering mit genetischen Algorithmen (GAs) in einer Cloud-Anwendung.

¹ Daimler TSS GmbH, Wilhelm-Runge-Straße 11, 89081 Ulm, natalia.ogulenko@daimler.com

² Daimler TSS GmbH, Wilhelm-Runge-Straße 11, 89081 Ulm, soeren.frey@daimler.com

³ Daimler TSS GmbH, Wilhelm-Runge-Straße 11, 89081 Ulm, jens.nahm@daimler.com

⁴ Hochschule Aalen - Technik und Wirtschaft, Wirtschaftsinformatik, Anton-Huber-Straße 25, 73430 Aalen, manfred.roessle@htw-aalen.de

1.1 Kontext

Viele Jahre waren Automobile fahrende "Daten-Inseln", d.h. sie waren in sich geschlossene Systeme. Seit ca. 20 Jahren werden die elektronisch gesteuerten Funktionseinheiten bei Werkstattbesuchen ausgelesen und die somit gesammelten Nutzungs- und Betriebsdaten weiterverwendet. Die seit ca. 10 Jahren eingesetzte Onlinevernetzung der Fahrzeuge wird als disruptiv bezeichnet, weil sie ganz neue Ansätze in vielen verschiedenen Dimensionen eröffnet.

Es steht beispielsweise eine viel größerer Umfang an Datenpunkten mit einer höheren Auflösung als bisher zur Verfügung. Die Erkenntnisse aus der Analyse dieser Daten fließen auch in die Entwicklung und Optimierung der nächsten Fahrzeuggeneration mit ein (Produktentwicklungszyklus ca. fünf Jahre). Darüber hinaus können Verbesserungen und die Beseitigung von Fehlern im Backend erfolgen und stehen dem Endnutzer sofort zur Verfügung (Produktentwicklungszyklus im Bereich von wenigen Tagen bis Wochen möglich).

In das Backend übermittelte Daten können etwa mit Methoden des maschinellen Lernens ausgewertet und mit aktuellen Kontextdaten wie z.B. zu Verkehr-, Wetter- und Straßenverhältnissen oder dem Luftschadstoffgehalt situativ per Onlineabruf aus dem Backend in das Fahrzeug übertragen werden.

Die rechen- und datenintensive Vorverarbeitung sprengt alle heute in Fahrzeugen verfügbaren Möglichkeiten. Die Onlinevernetzung der Fahrzeuge erlaubt verteilte Systemarchitekturen, bei denen die Speicherung der Daten und ein Großteil der Rechenlast in Cloud-basierte Backendsysteme verlagert wird. Durch die hohen Stückzahlen der produzierten Fahrzeuge zahlen sich selbst kleine Kostenoptimierungen pro Fahrzeug schnell aus. Wenn also anstatt eines sehr leistungsfähigen Rechners im Auto eine kleinere Variante ausreicht, die bei Bedarf mit Rechenleistung im Backend angereichert wird, können Kostenvorteile entstehen.

Ausgelagerte Datenanalysen im Backend eröffnen nun auch prinzipiell die Möglichkeit zur zentralen Optimierung von Fahrzeugbetriebsstrategien.

1.2 Problemstellung

Fahrzeugbetriebsstrategien zielen auf den effizienten Betrieb von Fahrzeugen ab (siehe Abschnitt 2.1) und werden bisher hauptsächlich statisch für eine Menge von Fahrzeugen, z.B. für bestimmte Baureihen, während deren Entwicklung optimiert. Das Potential zur dynamischen Optimierung einer Betriebsstrategie pro spezifischer, bevorstehender Fahrt durch eine zentrale Datenanalyse auf Basis von vielen historischen, ähnlichen Fahrten wird meistens noch nicht genutzt.

Eine entsprechende Datenanalyse könnte es z.B. ermöglichen, die energieeffizienteste Betriebsstrategie dynamisch für eine bevorstehende Fahrtroute für das jeweilige

Fahrzeug- und Motormodell und den gewählten Fahrstil (z.B. „Sport“), sowie unter Berücksichtigung der aktuellen Wetter- und Stauprognosen passend aus historischen Fahrtdaten mit ähnlichen Charakteristiken bestimmen zu können.

Formal stellt sich das adressierte Problem und das Ziel einer solchen Datenanalyse als die Minimierung einer Menge von Kriterien K dar. Im Kontext dieses Papiers sollen für eine geplante Fahrt mit einem Fahrzeug F von der Startlokation S (zum Zeitpunkt t_S) zur Ziellokation Z die Kriterien Kraftstoffverbrauch k_{KV} und Fahrzeit k_{FZ} ($K = \{k_{KV}, k_{FZ}\}$) minimiert werden. Hierbei soll eine beliebige Gewichtung von k_{KV} und k_{FZ} verwendbar sein. Für die Datenanalyse stehen hierzu folgende Informationen zu einer Menge von historischen Fahrten, sowie weitere Informationen zum jeweiligen Fahrkontext zur Verfügung:

- **Telemetriedaten:** Werden mehrfach pro Fahrt von Messinstrumenten im Fahrzeug aufgezeichnet und können an eine Cloud-Anwendung übermittelt werden. Z.B. Geschwindigkeit, Beschleunigung, Fahrstil (z.B. „Sport“), Geoposition.
- **Streckendaten:** Werden einer Cloud-Anwendung von externen Datenquellen bereitgestellt. Z.B. Geschwindigkeitsbegrenzungen oder Temperatur und Luftdruck bei einer bestimmten Geoposition und Uhrzeit.

Aufgrund der potenziell großen Menge an Telemetrie- und Streckendaten für viele verschiedene Fahrten muss ein geeignetes Verfahren zur Datenanalyse effizient berechenbar und skalierbar sein.

Viele allgemeine Optimierungsverfahren (z.B. Gradientenverfahren) setzen zur Minimierung von k_{KV} und k_{FZ} die Beschreibung des Zusammenhangs zwischen den Datenpunkten der Telemetrie-/Streckendaten und k_{KV} und k_{FZ} für bestimmte Fahrzeug- bzw. Motormodelle in Form von analytischen Modellen voraus. Diese Modelle sind jedoch z.B. für neue Datenpunkte der Telemetrie- und Streckendaten häufig nicht bekannt und müssen erst aufwendig erstellt werden. Um eine flexiblere Erweiterbarkeit um neue Fahrzeug- bzw. Motormodelle sowie neue Datenpunkte der Telemetrie- und Streckendaten zu erreichen, soll eine Datenanalyse zur Optimierung der Betriebsstrategie ohne entsprechende analytische Lösungsverfahren bzw. Simulationsmodelle verwendet werden können.

2 Hintergrund

2.1 Fahrzeugbetriebsstrategien

Fahrzeugbetriebsstrategien haben zum Ziel, ein Fahrzeug effizient betreiben zu können, z.B. hinsichtlich des Kraftstoffverbrauchs. Der Begriff *Fahrzeugbetriebsstrategie* (kurz *Betriebsstrategie*) wird zwar häufig im Kontext von Hybrid- und Elektrofahrzeugen verwendet, beschreibt aber allgemein die Auswahl von Betriebspunkten eines Fahrzeugs

(z.B. Gang bei Automatikgetrieben) [Ro14]. Die Auswahl erfolgt nach Kriterien wie z.B. der Beeinflussung der Energieeffizienz oder der Bauteilbelastung. Neben der direkten Steuerung von Kontrollsystemen im Fahrzeug können die berechneten Betriebspunkte einer Betriebsstrategie auch dem Fahrer zur manuellen Beeinflussung oder Kenntnisnahme signalisiert werden, indem z.B. eine Anzeige im Kombiinstrument oder ein haptisches Feedback am Gaspedal erfolgt.

Auch die Wahl einer Geschwindigkeit kann einen Betriebspunkt einer Fahrzeugbetriebsstrategie darstellen. Der Kraftstoffverbrauch kann durch die Wahl von geeigneten Geschwindigkeiten wesentlich beeinflusst werden [Le14]. Mit der Wahl einer adäquaten Geschwindigkeit kann z.B. der Einfluss des Luftwiderstands verringert und ein Verbrennungsmotor effizienter betrieben werden. Typischerweise steigt der Verbrauch bei geringen wie auch bei höheren Geschwindigkeiten ab einer bestimmten Grenze wieder an [Le14].

2.2 Genetische Algorithmen

Genetische Algorithmen (GAs) sind etablierte stochastische Suchverfahren zur Lösung von komplexen Optimierungsproblemen. Aufgrund ihrer Flexibilität kommen sie auch häufig bei Problemstellungen des Software Engineerings zum Einsatz (z.B. bei der Optimierung des Deployments einer Anwendung in einer Cloud Umgebung [FFH13]).

Die Funktionsweise von GAs ist dem Evolutionsprozess von Lebewesen nachempfunden [Po00]. GAs können auch Probleme adressieren, für deren Zielfunktion keine analytischen Modelle bekannt sind (im Gegensatz zu z.B. Gradientenverfahren, die differenzierbare Funktionen voraussetzen). GAs betrachten mehrere potenzielle Lösungskandidaten gleichzeitig. Diese Menge der Lösungen wird *Population* genannt. Die potenziellen Lösungen selbst werden als *Individuen* bezeichnet. Die einzelnen Bestandteile (z.B. Attribute) von Individuen werden *Gene* genannt. Alle Gene eines Individuums werden als *Chromosom* bezeichnet.

GAs arbeiten iterativ. In jeder Iteration wird eine neue *Generation* der Individuen erzeugt. Das Erzeugen von neuen Generationen wird mit Hilfe der Operatoren *Selektion*, *Rekombination* und *Mutation* durchgeführt. Diese Operatoren sind auch an natürliche Prozesse angelehnt und werden häufig spezifisch für bestimmte Problemdomänen erstellt.

Zuerst wird eine Menge an Individuen mit Hilfe des Selektions-Operators anhand der Bewertung mit einer Zielfunktion (sog. *Fitnessfunktion*) ausgewählt. Daraufhin wird der Rekombinations-Operator (auch *Crossover* genannt) zur Erzeugung von Nachkommen-Individuen verwendet, indem er Teile der Chromosome von Eltern-Individuen entsprechend bestimmter Regeln austauscht. Danach wird der Mutations-Operator angewandt, bei dem Teile eines Individuums zufällig verändert werden.

Parameter eines GAs, wie z.B. die Größe der Populationen oder die Anzahl an Generationen, werden *Kontrollparameter* genannt und können für ein gegebenes

Optimierungsproblem geeignet konfiguriert werden. GAs verursachen zwar oftmals einen erhöhten Rechenaufwand. Dieser Aspekt wird allerdings durch ihre sehr gute Parallelisierbarkeit abgeschwächt [Ha99].

2.3 Clustering mit Genetischen Algorithmen

Für die Bestimmung von optimalen Fahrzeugbetriebsstrategien verwenden wir Clustering auf Basis von GAs. *Clustering* (bzw. auch *Clusteranalyse* oder *Partitionierung*) bezeichnet hierbei Verfahren zur Aufteilung von Objekten in Klassen. Dabei sind die Kriterien für die Aufteilung vorab nicht bekannt. Somit gehört das Clustering zu der Kategorie des nicht-überwachten Lernens. Beim Clustering wird versucht, die Objekte so aufzuteilen, dass sie innerhalb eines Clusters möglichst ähnlich und in verschiedenen Clustern möglichst unterschiedlich sind [JMF99].

Häufig eingesetzte Verfahren wie K-Means haben oft den Nachteil, dass bei ihnen die Clusteranzahl der Ziel-Partitionierung vorgegeben werden muss. Diese Anzahl ist allerdings oftmals, wie auch in unserer Problemdomäne, nicht a priori bekannt. Diesem Nachteil wird häufig begegnet, indem das Clustering mehrfach mit unterschiedlichen Clusteranzahlen durchgeführt und die Verfahren somit kalibriert werden. Bei einer veränderlichen Datenbasis ist dabei eine regelmäßige Nachkalibrierung notwendig.

Clustering auf Basis von GAs [Gal1] adressiert die Vorgehensweise der Exploration des Lösungsraums mit verschiedenen Clusteranzahlen (in der von uns eingesetzten Variante) inhärent. Die Lösungskandidaten (Individuen) kodieren eine Menge von Clusterzentren, wobei die Clusteranzahl der Ziel-Partitionierung nicht vorgegeben werden muss. Da die Länge der Individuen nicht festgelegt ist, können während der Optimierung verschiedene Clusteranzahlen untersucht werden. Die Änderung der untersuchten Anzahl und Ausprägung der Clusterzentren wird durch die genetischen Operatoren Crossover und Mutation über mehrere Generationen hinweg durchgeführt.

Zur Bewertung der Lösungskandidaten, d.h. der Qualität der Partitionierung, verwenden wir in der Fitnessfunktion den etablierten *Davies-Bouldin-Index (DB-Index)* [Ar13]. Zur Bewertung der resultierenden endgültigen Clustering-Lösung wird von uns zusätzlich der *Silhouette-Index* [Ar13] eingesetzt, da er einen normalisierten Wert liefert, der auch allein (d.h. nicht nur im Vergleich zu anderen Partitionierungen) Auskunft über die Qualität eines Clusterings gibt. Da der DB-Index allerdings weniger Rechenressourcen als der Silhouette-Index bei gleichzeitig guten Ergebnissen liefert, ist die Anwendung des DB-Index in der Fitnessfunktion des GAs vorteilhaft.

3 Cloud-Basierte Optimierung von Betriebsstrategien

Das generelle Vorgehen zur Cloud-basierten Optimierung von Fahrzeugbetriebsstrategien wird in Abschnitt 3.1 beschrieben. Die Verwendung von Clustering zur Bestimmung von optimalen Betriebsstrategien in der Cloud erläutert

Abschnitt 3.2. Die speziellen Aspekte der GA-basierten Clusteranalyse beschreibt danach Abschnitt 3.3 im Detail.

3.1 Generelles Vorgehen

Die Optimierung der Betriebsstrategie für eine bevorstehende Fahrt macht sich die Erfahrungen aus vorhergehenden ähnlichen Fahrten unterschiedlicher Fahrer automatisiert zunutze. Abb. 1 zeigt einen Überblick über das generelle Vorgehen zur Cloud-basierten Optimierung von Fahrzeugbetriebsstrategien. Der Fahrer gibt im Fahrzeug das gewünschte Reiseziel ein (z.B. in dessen Infotainmentsystem) und stellt darüber hinaus den präferierten Fahrstil ein bzw. verwendet eine Voreinstellung. Daraufhin wird die Route berechnet und zusammen mit weiteren Informationen, wie z.B. zum gewählten Fahrstil oder dem verwendeten Fahrzeug(modell), zur Cloud-Anwendung übertragen. Dort findet kontinuierlich eine Aggregation der Telemetriedaten aller übermittelten Fahrten, sowie der zugehörigen Streckendaten statt (weitere Details zu dieser Datenaggregation werden in Abschnitt 3.2 erläutert).

Mittels eines Clusterings dieser Daten auf Basis eines GAs wird der Suchraum für Fahrten mit ähnlichen Charakteristiken stark eingeschränkt. Somit lässt sich für eine bevorstehende Fahrt der Cluster mit den ähnlichsten Telemetrie- und Streckendaten auswählen und die Suche nach der jeweiligen optimalen Betriebsstrategie auf diesen Cluster begrenzen. Die Realisierung der Clusteranalyse mit Hilfe eines GAs unterstützt hierbei auch das Ziel einer guten Skalierbarkeit (vgl. Abschnitt 1.2).

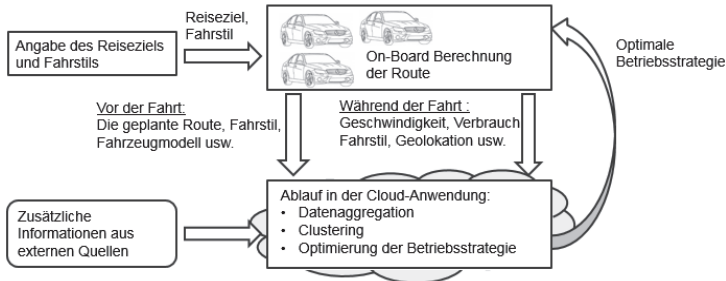


Abb. 1: Generelles Vorgehen zur Cloud-basierten Optimierung von Betriebsstrategien

Durch eine vorausschauende Fahrstrategie soll eine Minimierung des Kraftstoffverbrauchs und der Fahrzeit (k_{KV} und k_{FZ} , siehe Abschnitt 1.2) erreicht werden. Die Regulierung von k_{KV} und k_{FZ} erfolgt über die Berechnung einer idealen Geschwindigkeit (vgl. Abschnitt 2.1) für einzelne Streckenabschnitte in einem gegebenen Kontext (z.B. Verkehrsaufkommen, Wetterbedingungen, Motorleistung).

Eine optimale Betriebsstrategie für eine Route setzt sich aus den optimalen berechneten Geschwindigkeiten für die einzelnen Streckenabschnitte der Route zusammen. Nachdem die optimale Betriebsstrategie für eine Route in der Cloud-Anwendung berechnet wurde,

kann diese in das Fahrzeug übermittelt und dort angewandt werden. Während der Fahrt werden aktuelle Telemetrie- und Streckendaten kontinuierlich zur Cloud-Anwendung übertragen. Somit kann der Vorgang zur Bestimmung und Übermittlung der optimalen Betriebsstrategie während der Fahrt wiederholt werden, um beispielsweise auf veränderte Stau- oder Wetterprognosen reagieren zu können.

3.2 Bestimmung von Optimalen Betriebsstrategien mittels Clustering

Der Ansatz zur Bestimmung von optimalen Betriebsstrategien in einer Cloud-Anwendung unter Verwendung von Clustering ist in Abb. 2 dargestellt. In der Cloud-Anwendung findet die Analyse der übermittelten Daten mittels Clustering statt. Hier wird zuerst die Reiseroute und ihr Kontext, wie z.B. die Antriebsart des Fahrzeugs, Temperatur und für die Tageszeit charakteristische Stauprognose, in Abschnitte von 100 Metern aufgeteilt (Schritt 1). Hierbei werden die durchschnittlichen Werte pro Abschnitt berechnet. Für jeden Abschnitt wird der Cluster mit den ähnlichsten Strecken bestimmt (gemessen am euklidischen Abstand eines Abschnitts zu einem Clusterzentrum) und der jeweilige Abschnitt wird diesem Cluster zugeordnet (Schritt 2). Für jeden dieser Cluster wird das erste Prozent der Abschnitte aus den historischen Daten bestimmt, die die höchste Ähnlichkeit (auch gemäß dem euklidischen Abstand) mit dem zu analysierenden Abschnitt aufweisen (Schritt 3). Aus der Menge der ähnlichen Abschnitte werden diejenigen ausgewählt, die die am besten geeignete Betriebsstrategie aufweisen (Schritt 4).

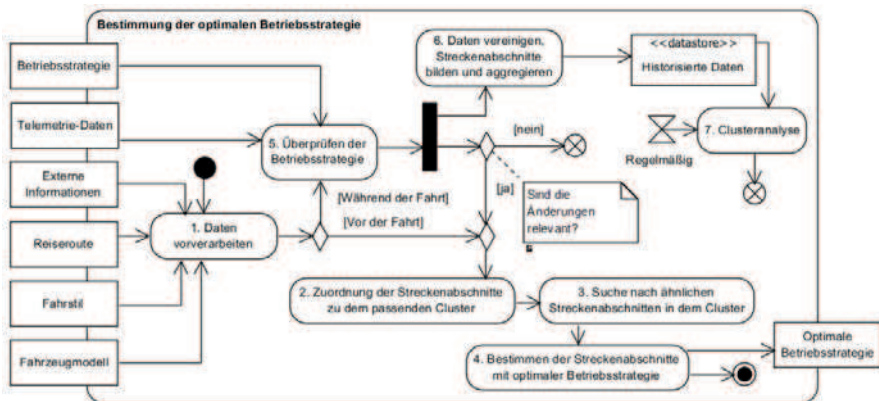


Abb. 2: Ansatz zur Bestimmung der optimalen Betriebsstrategie mittels Clustering

Hierbei wird der Abschnitt mit der geringsten normalisierten Summe von Kraftstoff- bzw. Energieverbrauch und Fahrzeit, entsprechend ihren Gewichten, ausgewählt. Da jedem Streckenabschnitt eine durchschnittliche Geschwindigkeit zugeordnet ist, wird auf diese Weise die optimale Geschwindigkeit für jeden geplanten Streckenabschnitt ermittelt. Anschließend kann die gesamte optimale Betriebsstrategie für die ganze

bevorstehende Strecke in das Fahrzeug übermittelt werden. Eine Priorisierung von Zeit und Kraftstoff- bzw. Energieverbrauch kann vom Fahrer gemäß seiner Präferenzen variiert werden.

Wenn sich die Bedingungen während der Fahrt ändern, wird die alte Betriebsstrategie überprüft (Schritt 5). Falls die Änderungen relevant sind, z.B. bei veränderter Stauprognose oder bei der Änderung der Reiseroute, wird eine neue optimale Betriebsstrategie durch erneutes Ausführen der Schritte 2-4 bestimmt. Die neue Betriebsstrategie ersetzt in diesem Fall die alte.

Die Datenbasis für die Optimierungen wird aus historischen Daten geschaffen. Während der Fahrt werden vom Fahrzeug erzeugte Telemetrie-Daten kontinuierlich in die Cloud-Anwendung übertragen. Im nächsten Schritt werden die Fahrzeugdaten in 100 Meter langen Streckenabschnitten aggregiert und in der Cloud-Anwendung abgelegt (Schritt 6). Im Unterschied zu Schritt 1 werden hierbei auch für jeden Streckenabschnitt die durchschnittlichen Werte der dynamischen Sensormesswerte (u.a. auch der Geschwindigkeit) berechnet. Die in Abschnitten von 100 Metern aggregierten Daten werden mit zusätzlichen Informationen, z.B. zu Fahrzeugcharakteristiken, Wetterbedingungen, Verkehrsaufkommen und dem eingestellten Fahrstil, ergänzt. Anschließend wird die Clusteranalyse durchgeführt (Schritt 7), d.h. die Aufteilung von Streckenabschnitten mit ähnlichen Eigenschaften in Cluster zur Einschränkung des Suchraums. Da die neuen Fahrzeug- und Streckendaten (ohne Bezug zu Fahrerinformationen) kontinuierlich zu der Datenbasis hinzugefügt werden, wird auch das Clustering regelmäßig durchgeführt. Die GA-basierte Clusteranalyse wird in Abschnitt 3.3 detailliert beschrieben. Generell können die Länge der Streckenabschnitte und der Anteil der ähnlichsten Abschnitte in Abhängigkeit von den zur Verfügung gestellten Rechenkapazitäten variiert werden.

3.3 GA-Basierte Clusteranalyse

Bei der Clusteranalyse mit GA sind die potenziellen Lösungen (eine Lösung ist eine Aufteilung in Cluster) wie folgt aufgebaut: Die Gene der Chromosome kodieren Clusterzentren. Bei n Attributen der Telemetrie- und Streckendaten repräsentiert jedes Clusterzentrum pro Streckenabschnitt einen Punkt im n -dimensionalen Raum (Abb. 3). Durch diese Repräsentation lassen sich auch neue Datenattribute der Telemetrie- und Streckendaten leicht in die Clusteranalyse integrieren.

Die Länge der Chromosome variiert im Intervall zwischen minimaler und maximaler Anzahl der Clusterzentren. Das Intervall wird für den Kontext der experimentellen Evaluation (vgl. Abschnitt 4) von 2 bis 100 Clusterzentren empirisch festgelegt, es sind jedoch auch Verfahren zur dynamischen Adaption des Intervalls denkbar. Bei der Erzeugung der ersten Population wird die Anzahl der Clusterzentren für jedes Individuum (d.h. jedes Chromosom) zufällig aus diesem Intervall ausgewählt.

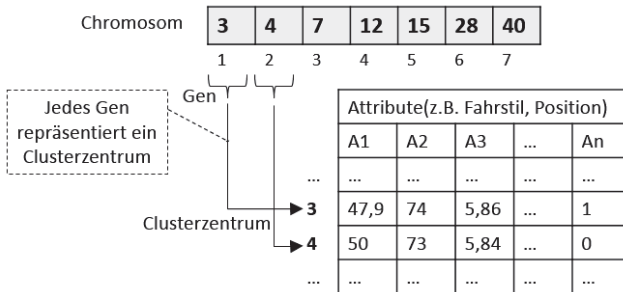


Abb. 3: Ein exemplarisches Chromosom

Die Berechnung des Fitnesswerts von Individuen wird mit einer Fitnessfunktion realisiert, bei der zuerst alle Streckenabschnitte der historischen Daten anhand des niedrigsten euklidischen Abstands zu bestehenden Clusterzentren des Individuums zugeordnet werden. Daraufhin wird die Aufteilung mit dem DB-Index (vgl. Abschnitt 2.3) bewertet. Als *Selektions*-Operator wird eine Fitness-proportionale Selektion mit linearer Skalierung angewendet. Die Selektion erfolgt auf Basis von Auftrittswahrscheinlichkeiten, die den Individuen entsprechend ihrer Fitnesswerte zugeordnet werden. Um eine vorzeitige Konvergenz auf ein lokales Optimum zu vermeiden, wird bei der Selektion ein Rauschen [Po00] hinzugefügt. Hierdurch werden auch ein Prozent der Individuen mit unterdurchschnittlichen Fitnesswerten für die Erzeugung der nächsten Generation verwendet.

Die Reihenfolge der Clusterzentren in Individuen ist bei der *Rekombination (Crossover)* unwichtig. Von Bedeutung ist die Abwesenheit von Duplikaten in jedem Nachkommen-Individuum (Kind). Deswegen werden die einzigartigen Gene, d.h. die Clusterzentren, der zwei Eltern-Individuen bei der Rekombination in einen Pool abgelegt. Aus diesen Genen werden zwei Nachkommen erzeugt. Die Anzahl ihrer Clusterzentren wird wie folgt aus den Clusterzentren ihrer Eltern berechnet. Wenn $k1$ und $k2$ die Länge der zwei Eltern-Individuen und $h1 (= \lceil k1/2 \rceil)$ und $h2 (= \lfloor k2/2 \rfloor)$ Kreuzungspunkte (bzw. die Länge der Eltern-Individuen bis zu diesen Kreuzungspunkten) bezeichnen, wird die Länge der Nachkommen-Individuen als $h1+k2-h2$ und $h2+k1-h1$ festgelegt. Nach der Erzeugung des ersten Kindes (durch das zufällige Entnehmen von Genen aus dem Pool) werden die in dem Genpool noch nicht verwendeten Gene für die Erzeugung des zweiten Kindes zufällig entnommen. Falls diese verbleibenden Gene nicht ausreichen, werden für die Erzeugung des restlichen Teils des zweiten Kindes gleiche (zufällig ausgewählte) Gene wie bei dem ersten Kind verwendet. Abb. 4 zeigt dieses Vorgehen für den Rekombinations-Operator an einem Beispiel, bei dem das sechste Gen des ersten und fünfte Gen des zweiten Eltern- Individuums gleich sind. Somit wird das Gen nur einmal in den Genpool übernommen.

Bei der *Mutation* wird ein Clusterzentrum in einem Individuum durch ein anderes zufällig ersetzt und die Clusteranzahl zufällig vergrößert oder verkleinert.

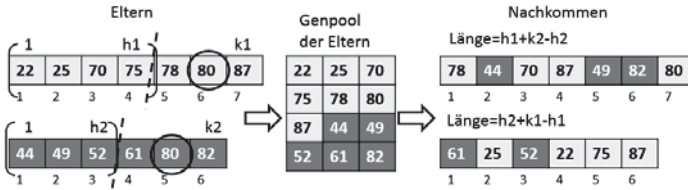


Abb. 4: Rekombination und Vorgehen bei doppelten Clusterzentren (markiert mit Kreis)

Hierbei werden die Individuen, die weniger als sechs Clusterzentren enthalten, nicht verkleinert. Diese empirisch bestimmte Grenze verhindert effektiv die Konvergenz auf nur ein Clusterzentrum. Die Anzahl der Clusterzentren, die abgeschnitten oder hinzugefügt werden, wird zufällig aus dem Intervall von Null bis zur Hälfte der Anzahl der Clusterzentren des ursprünglichen Individuums ausgewählt. Bei der Mutation wird darauf geachtet, dass zu einem Individuum nur Clusterzentren hinzugefügt werden, die es noch nicht enthält. In Abb. 5 wird das Vorgehen bei der Mutation illustriert. Das zu mutierende Individuum umfasst sieben Gene. Im ersten Schritt wird zufällig das Gen an fünfter Position mit einem anderem (dunkelgrau markiert) ersetzt. Im zweiten Schritt wird entschieden, ob das Individuum verkürzt oder verlängert werden muss. In diesem Beispiel werden die letzten beiden Gene abgetrennt.

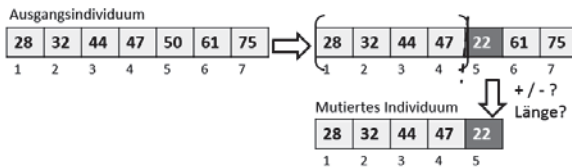


Abb. 5: Mutations-Operator

4 Experimentelle Evaluation

Als Datenbasis werden aggregierte historische Daten aus 62 aufgezeichneten Fahrten verwendet, die mit drei unterschiedlichen Fahrzeugen mit verschiedenen Antriebsarten und in unterschiedlichen Regionen für die Evaluation durchgeführt wurden. Nach der Aggregation der Daten zu Abschnitten von 100 Metern entstanden hierbei 14.449 Datensätze. Bei der Datenanalyse werden für die Evaluation 17 Datenattribute verwendet, u.a. die Steigung und Kurvigkeit der Streckenabschnitte, die Luftfeuchtigkeit und Umgebungstemperatur, der eingestellte Fahrmodus, sowie das vom Hersteller angegebene maximale Drehmoment und die Leistung und Beschleunigung eines Fahrzeug- bzw. Motormodells. Die Abtastrate der Telemetriedaten beträgt 2 Hz, d.h. alle 0,5 Sekunden werden die Telemetriedaten ausgelesen und gespeichert.

4.1 Validierung des Clusterings

Um die Kontrollparameter des GA empirisch zu bestimmen, die in unserem Kontext die besten Fitnesswerte ermöglichen, werden Versuche mit variierenden Kontrollparametern durchgeführt. Die folgenden Parameter erweisen sich in Bezug auf den DB-Index und den Silhouette-Index als am besten geeignet [Og16]. Die Populationsgröße beträgt lediglich 100 Individuen. Deswegen wird eine große Anzahl von Generationen (5.000) festgelegt. Der Algorithmus terminiert, wenn der beste Fitnesswert während 300 Iterationen konstant bleibt. Eine hohe Crossoverrate (80%) und eine relativ kleine Mutationsrate (10%) werden in Kombination mit Elitismus [Po00] (4%) pro Generation angewandt.

Mit dieser Konfiguration findet der GA in den Fahrtdaten drei Cluster. Zur Bewertung der Ergebnisse des Clusterings wird der Silhouette-Index verwendet, dieser liegt generell im Intervall $[-1;1]$, wobei größere Werte bessere Ergebnisse repräsentieren und positive Werte dem Fall entsprechen, dass die Abstände zu den anderen Clustern größer sind als die Distanz innerhalb der Cluster. Der Silhouette-Index der berechneten Clusterlösung beträgt $\sim 0,48$. Dies stellt somit ein gutes Ergebnis dar.

In Abb. 6 werden exemplarisch die Verteilungen der Attribute „Pedalposition“ (links, wie weit ist ein Gaspedal gedrückt) und „Lufttemperatur“ (rechts) pro Streckenabschnitt in jedem Cluster dargestellt. Cluster 1 separiert z.B. deutlich die Abschnitte, bei denen mehr Gas gegeben wurde und solche mit niedrigen Temperaturen.

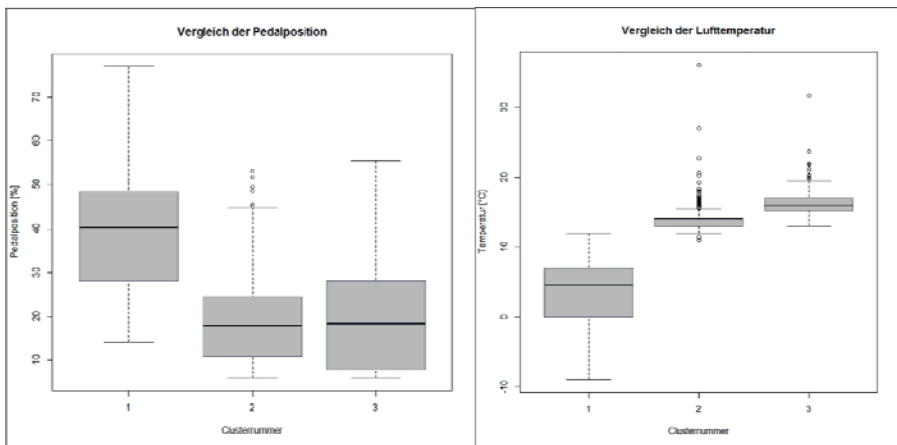


Abb. 6: Verteilung von Pedalposition (links) und Lufttemperatur (rechts) in den Clustern

Um den Unterschied in den Ergebnissen des Clusterings in Abhängigkeit von der Größe der Datenbasis festzustellen, wird der GA mit verschiedenen großen Anteilen der Datenbasis evaluiert. Hierbei wird auch die sequenzielle Ausführung des GAs mit der parallelen Ausführung verglichen (16 GB RAM, CPU mit 4 Kernen, Intel® Core(TM) i5-2400 CPU @ 3.10 GHz). Die Ergebnisse sind in Abb. 7 dargestellt.

Die linke Grafik in Abb. 7 veranschaulicht die ansteigende Rechenzeit des Algorithmus pro Iteration mit dem Anstieg des Datenvolumens. Jedoch bringt die parallele Berechnung des Algorithmus einen deutlichen Performancegewinn mit sich. Die Rechenzeit verringert sich über alle Datenvolumina durchschnittlich um ca. die Hälfte. Dies entspricht den Erwartungen, da bei paralleler Verarbeitung der Performancegewinn nicht linear mit der Zahl der Cores skaliert [PKG16]. Die rechte Grafik in Abb. 7 zeigt die Abhängigkeit des Fitnesswertes vom Datenvolumen. Der maximale Fitnesswert wird mit 100% der Daten erreicht und unterscheidet sich bei der sequenziellen und parallelen Ausführung des GAs fast nicht.

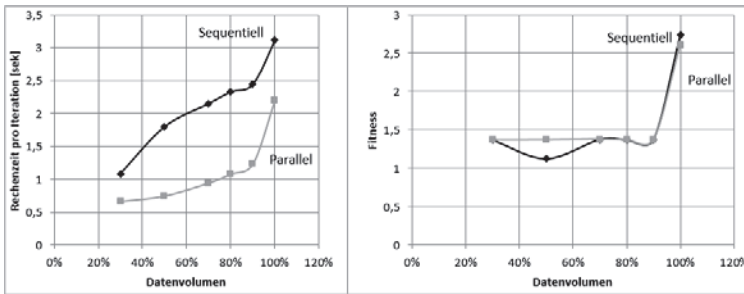


Abb. 7: Vergleich der sequenziellen mit der parallelen Ausführung des GAs

4.2 Bestimmung der optimalen Betriebsstrategien

Für die Validierung der Bestimmung der optimalen Betriebsstrategien werden mit unserem Ansatz die optimalen Geschwindigkeiten für zwei Strecken berechnet und mit den tatsächlichen Geschwindigkeiten während der Fahrt verglichen. Hierbei werden die Fahrzeit, der Spritverbrauch und die CO₂-Emissionen pro Strecke für unterschiedliche Optimierungskriterien berechnet, d.h. für unterschiedliche Gewichtungen der Zeit und des Spritverbrauchs. Den Berechnungen zufolge sinkt der Kraftstoffverbrauch der Fahrzeuge für beide Strecken mit der berechneten optimalen Betriebsstrategie. Falls das Optimierungskriterium zu 100% den Spritverbrauch priorisiert, sind, wie erwartet, die CO₂-Emissionen und der Spritverbrauch pro Strecke am geringsten. Falls das Optimierungskriterium zu 100% die Zeit priorisiert, ist die Zeit pro Strecke am kürzesten, die CO₂-Emissionen und der Spritverbrauch erhöhen sich jedoch.

5 Verwandte Arbeiten

Terwen [Te09] beschäftigt sich mit der vorausschauenden Längsregelung schwerer Lastkraftwagen und Radke [Ra13] mit der energieoptimalen Längsführung von Personenkraftwagen. Beide Arbeiten zielen auf die Reduktion des Kraftstoffverbrauchs durch den Einsatz vorausschauender Fahrstrategien. Die Optimierung wird bei Terwen [Te09] mit Hilfe der Quadratischen Programmierung für modellbasierte prädiktive

Regelung und bei Radke mittels Dynamischer Programmierung in Echtzeit im Fahrzeug durchgeführt. Der Gang und wie bei unserem Ansatz die Fahrgeschwindigkeit sind die Basis für die Bestimmung der optimalen Betriebsstrategie [Ra13]. Genaue Angaben zum Umfang der Betriebspunktendatenbanken wie beispielsweise die Zahl der gespeicherten Streckenabschnitte finden sich nicht.

Unter Verwendung der Dynamischen Programmierung wurde von Roth [Ro14] ein Optimierungsansatz auf Basis einer Betriebspunktendatenbank vorgestellt. Hierbei werden die per Simulation erzeugten Fahrzeug-Betriebspunkte auf Streckenabschnitten von fünf Metern aggregiert und in einer Datenbank gespeichert.

Die Arbeit von Salmasi [Sa07] beschäftigt sich mit Regelungsstrategien für Hybrid-Fahrzeuge. Er verwendet zur Darstellung die vier Klassen „Deterministische“ und „Fuzzy“ regelbasierte Ansätze sowie „Globale“ und „Echtzeit“ Optimierungsansätze und stellt die jeweiligen Verfahren mit ihren mathematischen Grundlagen dar. Ein konkreter Anwendungsfall findet sich jedoch nicht.

Zusammenfassend bleibt festzuhalten, dass alle hier referierten Arbeiten im Wesentlichen dezentrale, fahrzeugbasierte Ansätze verfolgen. Unser Ansatz nutzt genetische Algorithmen wie sie auch in [Sa07] dargestellt sind, geht jedoch durch die Nutzung einer Cloudlösung einen wichtigen Schritt hin zur fahrzeugunabhängigen Nutzung der Streckendaten und Optimierungsansätze.

6 Zusammenfassung und Ausblick

Die Optimierung von Fahrzeugbetriebsstrategien erfolgte bisher hauptsächlich statisch für eine Menge von Fahrzeugen, z.B. für bestimmte Baureihen, während deren Entwicklung. Dieses Papier beschreibt einen Ansatz, der sich die zunehmende Onlinevernetzung von Fahrzeugen zunutze macht und Telemetrie- und Streckendaten von vielen unterschiedlichen Fahrten vieler unterschiedlicher Fahrzeuge in eine Cloud-Anwendung überträgt. Mittels eines Clusteringverfahrens auf Basis von genetischen Algorithmen können daraufhin für bevorstehende Fahrten optimale Betriebsstrategien berechnet werden, die für vorliegende spezifische Charakteristika (z.B. Verkehrsaufkommen, Temperatur, Motorleistung) den Kraftstoffverbrauch und/oder die Fahrtzeit minimieren.

Für zukünftige Erweiterungen unseres Ansatzes sind besonders die Einbeziehung von weiteren Attributen der Telemetrie- und Streckendaten relevant. Des Weiteren könnten durch die Variation der Auflösung der übermittelten Daten, sowie der Länge der Streckenabschnitte, die für die Aggregation dieser Daten verwendet werden, weitere Qualitäts- und Performanceverbesserungen erzielt werden.

Literaturverzeichnis

- [Ar13] Arbelaitz, O. et. al.: An extensive comparative study of cluster validity indices. In: Pattern Recognition, Vol. 46, Issue 1, Elsevier, S. 243–256, 2013.
- [FFH13] Frey, S.; Fittkau, F.; Hasselbring, W.: Search-Based Genetic Optimization for Deployment and Reconfiguration of Software in the Cloud. In: Proceedings of the 35th International Conference on Software Engineering (ICSE 2013), IEEE Press, S. 512–521, 2013.
- [Ga11] Gajawada, S. et. al.: Optimal Clustering Method Based on Genetic Algorithm. In: Proceedings of the International Conference on Soft Computing for Problem Solving (SocProS 2011), 2011.
- [Ha99] Hansohm, J.: Clusteranalyse mit Genetischen Algorithmen. In: Mathematische Methoden der Wirtschaftswissenschaften, Springer, S. 57–66, 1999.
- [JMF99] Jain, A. K.; Murty, M. N.; Flynn, P. J.: Data Clustering: A Review, ACM Computing Surveys, Vol. 31, Issue 3, S. 264–323, 1999.
- [KW14] Köhler, T. R.; Wollschläger, D.: Die digitale Transformation des Automobils: 5 Mega-Trends verändern die Branche. Media-Manufaktur, 2014.
- [Le14] Lederer, M.: Energieeffizientes Fahren. In: Energiemanagement im Kraftfahrzeug. Optimierung von CO₂-Emissionen und Verbrauch konventioneller und elektrifizierter Automobile, Springer Vieweg, S. 307–321, 2014.
- [Og16] Ogulenko, N.: Cloud-basierte Optimierung von Fahrzeugbetriebsstrategien auf Basis historischer Telemetrie- und Streckendaten, Masterthesis, Hochschule Aachen - Technik und Wirtschaft, 2016.
- [PKG16] Pei, S.; Kim, M.-S.; Gaudiot, J.-L.: Extending Amdahl's Law for Heterogeneous Multicore Processor with Consideration of the Overhead of Data Preparation. In: IEEE Embedded Systems Letters, Vol. 8, No. 1, S. 26–29, 2016.
- [Po00] Pohlheim, H.: Evolutionäre Algorithmen - Verfahren, Operatoren und Hinweise für die Praxis, Springer, 2000.
- [Ra13] Radke, T.: Energieoptimale Längsführung von Kraftfahrzeugen durch Einsatz vorausschauender Fahrstrategien. Karlsruher Institut für Technologie (KIT) Fakultät für Maschinenbau: KIT Scientific Publishing, Dissertation, 2013.
- [Ro14] Roth, M.: Betriebsstrategie. In: Energiemanagement im Kraftfahrzeug. Optimierung von CO₂-Emissionen und Verbrauch konventioneller und elektrifizierter Automobile, Springer Vieweg, S. 323–365, 2014.
- [Sa07] Salmasi, F.: Control Strategies for Hybrid Electric Vehicles: Evolution, Classification, Comparison, and Future Trends. In: IEEE Transactions on Vehicular Technology, S. 2393–2404, 2007.
- [Te09] Terwen, S.: Vorausschauende Längsregelung schwerer Lastkraftwagen. Schriften des Instituts für Regelungs- und Steuerungssysteme Karlsruher Institut für Technologie: KIT Scientific Publishing, Dissertation, 2009.