

Unüberwachtes Lernen von KI-Systemen bei der Auswertung von landwirtschaftlichen Prozessen

Thoralf Stein¹

Abstract: In dieser Arbeit sollen Einsatzmöglichkeiten von KI-Systemen gezeigt werden, die auf dem sogenannten unüberwachten Lernen basieren und somit keine Datengrundlage und kein Training benötigen. Dabei werden verschiedene Einsatzgebiete rund um den Maschineneinsatz und dessen Auswertung in der Landwirtschaft gezeigt. Deren Stärke liegt in dem Gruppieren von Daten und in dem Finden von Ausreißern. Es werden Vor- und Nachteile der Algorithmen gezeigt und der tatsächliche Nutzen im Einsatz diskutiert.

Keywords: unüberwachtes Lernen, DBSCAN, k-means, landwirtschaftliche Prozesse, Datenauswertung

1 Einleitung

Bereits viele Anwendungen in der Landwirtschaft nutzen künstliche Intelligenz (KI) in den verschiedensten Formen. Auch in Bezug auf mobile Arbeitsmaschinen bieten sich diverse Möglichkeiten, KI-Systeme einzusetzen. Bekannte Anwendungen sind beispielsweise Erkennung von Arbeitszuständen [Br21], Bestimmung des Betriebsmodus [Po20] sowie Erkennung von Anbaugeräten [SM18]. Alle diese Ansätze basieren auf dem überwachten Lernen (engl. „supervised machine learning“). Das bedeutet, dass bereits eine Datenbasis mit Beobachtungen und Ergebnissen für das Lernen der KI vorhanden sein muss, sodass die KI mit diesen Daten trainiert werden kann und anschließend einsatzbereit ist.

Es gibt auch diverse Anwendungen in der Landwirtschaftsinformatik, die auf dem unüberwachten Lernen (engl. „unsupervised machine learning“) basieren, jedoch bezieht sich ein Großteil der Arbeiten auf pflanzenbauliche Analysen sowie die Bilderkennung in zahlreichen Anwendungen. In diesem Beitrag wird gezeigt, wie solche KI-Systeme im Zusammenhang mit landwirtschaftlichen Prozessen und der dazugehörigen Datenaufzeichnung sowie -Auswertung genutzt werden können. Dargelegt werden die Einsatzmöglichkeiten anhand von Daten, die im Rahmen der Projekte BiDa-LAP sowie BiDa-LAP II aufgezeichnet worden sind. Es wurden drei Testbetriebe mit Kommunikationsmodulen ausgestattet, die kontinuierlich Motor- und GPS-Daten der Arbeitsmaschinen in einem 1-Hz-Takt sammeln. Dabei werden die Module von dem Projektpartner Logicway GmbH bereitgestellt, der Praxisbetrieb wird von der Technischen Universität Dresden, Fachgebiet

¹ Technische Universität Berlin, Konstruktion von Maschinensystemen, Straße des 17. Juni 144, 10623 Berlin, thoralf.stein@tu-berlin.de

Agrarsystemtechnik bewerkstelligt und die Serverarchitektur sowie der Webservice von der Agricon GmbH bereitgestellt. Verschiedene Auswertungsalgorithmen in Kombination mit Big-Data-Techniken werden für das Projekt von der Technischen Universität Berlin, Fachgebiet Konstruktion von Maschinensystemen, entwickelt.

In dem Beitrag werden verschiedene Algorithmen vorgestellt, die auf dem unüberwachten Lernen basieren. Gezeigte Anwendungen sind die Bewertung der Qualität der aufgezeichneten Daten, optimierte Datenauswertung durch vereinfachtes Bestimmen der Zeitgliederung der Arbeiten sowie das Finden von Schlägen ohne zusätzliche Daten.

2 Unüberwachtes Lernen

Die Einsatzmöglichkeiten von Systemen, die auf dem unüberwachten Lernen basieren, sind weit gefächert. In der Auswertung von landwirtschaftlichen Prozessen können die Systeme helfen, die Analyse zu erleichtern sowie möglicherweise neue Einsichten oder unentdeckte Effekte zu offenbaren.

Sie können auch für die Komprimierung beziehungsweise für die Verringerung der Dimension von Datensätzen genutzt werden, um deren Handhabung zu vereinfachen. Die am häufigsten genutzten Algorithmen sind dabei die Hauptkomponentenanalyse (HKA) [Jo10] sowie der Autoencoder [Pi12]. Beide können Datensätze deutlich vereinfachen, die HKA basiert dabei auf der Berechnung von Eigenvektoren und Kovarianzmatrix, der Autoencoder arbeitet mit neuronalen Netzen. Genutzt werden beide oftmals in Kombination mit überwachtem Lernen, um den Rechenaufwand zu minimieren.

Eine weitere große Stärke von unüberwachten KI-Systemen liegt im eigenständigen Segmentieren und Gruppieren von Datensätzen [Jo21]. Dabei werden nach bestimmten Kriterien Objekte mit ähnlichen Eigenschaften zu Gruppen zusammengefasst. Wie genau dabei vorgegangen wird, hängt von der Zielsetzung bei der Gruppierung sowie von dem genutzten Algorithmus ab. Auch die Qualität der Daten sowie die benötigte Rechenzeit spielen eine Rolle für das Vorgehen.

In Abbildung 1 wurden beispielhaft 2-dimensionale Daten, die zum Beispiel räumliche Daten sein können, durch einen k-means-Algorithmus gruppiert. Dieser basiert darauf, Daten zu einer vorgegeben Anzahl Gruppen einzuteilen und die Daten abhängig von der Varianz zum Gruppenmittelpunkt zuzuordnen, woher auch die Namensgebung stammt [Ma03]. Anschließend werden anhand der gebildeten Gruppen die Mittelpunkte neu berechnet und dieses Vorgehen wird so lange wiederholt, bis sich die Mittelpunkte nicht mehr verschieben. Der Algorithmus kann nicht nur mit räumlichen Daten arbeiten, sondern liefert auch stabile Ergebnisse bei mehrdimensionalen Daten, die Eigenschaften von Objekten entsprechen können.

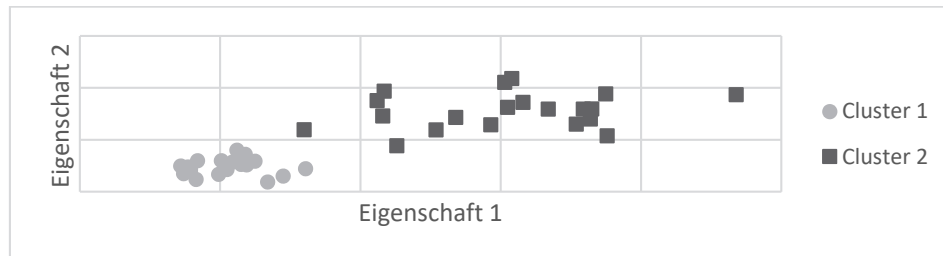


Abb. 1: Beispielhafte Darstellung von k-means-gruppierten 2-D Daten

Weitere Algorithmen werden im Folgenden erläutert und deren Anwendung beschrieben. Alle im Folgenden gezeigten Ansätze wurden mittels dem Programm Matlab der Firma MathWorks entwickelt und getestet. An dieser Stelle wird darauf hingewiesen, dass die Algorithmen nur ansatzweise erklärt werden können, andernfalls würde dies den Rahmen des Beitrags übersteigen.

3 Dichte-basierter Scan

Der Dichte-basierte Scan (DBSCAN) hat den Vorteil, nicht nach einer bestimmten Anzahl von Clustern zu gruppieren, sondern es wird nach der „Dichte“ der Daten vorgegangen. Dabei muss nicht zwangsläufig jeder Datenpunkt einer Gruppe zugeordnet werden, sondern es werden erst dann Cluster erstellt, wenn bestimmte Grenzwerte erreicht werden. Diese Grenzwerte sind zum einen Epsilon, was bei 2D-Daten einem Radius entspricht, sowie eine minimale Anzahl von Punkten, aus denen ein Cluster bestehen muss [Es96]. Bei der Auswertung von landwirtschaftlichen Prozessen kann der Algorithmus für verschiedene Anwendungen benutzt werden. Zum einen können Ausreißer in den Daten gefunden werden, die auf Anomalien in der Aufzeichnung hindeuten sowie mögliche Probleme des GPS-Signals bedeuten können. Zwar liefern Datenlogger oftmals auch Qualitätsdaten wie die Anzahl der Satelliten sowie HDOP oder PDOP mit, aber nicht immer lassen sich daraus Probleme erkennen, und DBSCAN kann in diesem Fall unterstützend genutzt werden.

Eine weitere Anwendungsmöglichkeit liegt in der Zuweisung von Schlägen. Dabei können die GPS-Daten nach der Dichte und Anzahl gruppiert werden und es lassen sich gut die Schläge ausfindig machen. Auch dies kann unterstützend genutzt werden, falls die Schlagkartei unvollständig oder nicht vorhanden ist. Für beide Anwendungen müssen die Eingangsparameter Epsilon sowie die minimale Punktzahl bestimmt werden. Zur Detektion von Anomalien kann man den Algorithmus an eine Schleife koppeln, die stoppt, sobald erste Cluster gebildet werden, so können Anomalien in den aufgezeichneten Daten gefunden werden. Das Generieren der Schläge erfolgt mittels eines Epsilon, welches etwa dem Dreifachen der Spurweite entspricht, die ebenfalls automatisch bestimmt werden kann. Eine minimale Punktzahl von 60 hat sich für die Erstellung der Schläge als robust erwiesen. In Abbildung 2 wurde dies beispielhaft dargestellt.

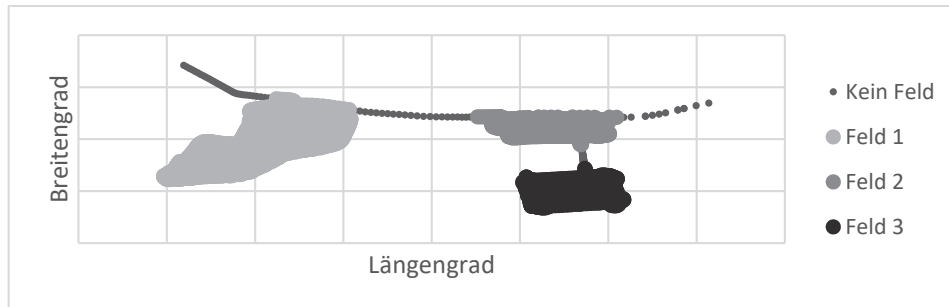


Abb. 2: GPS-Daten mittels DBSCAN gruppiert, Straßen werden keiner Gruppe zugeordnet

Dabei ist die Qualität des Ergebnisses stark abhängig von der Art der Arbeit, optimale Ergebnisse wurden bei der Bodenbearbeitung erzielt. 1200 Aufzeichnungstage wurden analysiert und dabei wurden die Ergebnisse des DBSCAN mit den tatsächlichen Feldgrenzen abgeglichen und es hat sich eine durchschnittliche Übereinstimmung von 75 % ergeben, was als gutes Ergebnis zu bewerten ist, da unterschiedlichste Arbeiten mit einbezogen wurden. Andere Arbeiten, die die Bestimmung nach Feld oder Straße vornehmen [PEN20], benutzen überwachtes Lernen und haben dadurch eine Trefferquote von über 92 %, bringen aber den Aufwand des Trainings sowie der Validierung mit sich. Weiterhin gibt es Arbeiten, die eine Erkennung mittels eines analytischen Verfahrens nutzen, wie etwa [He14]. Dazu sind jedoch keine Gütemaße bekannt.

4 K-Means-Cluster

Eine weitere mögliche Anwendung eines Algorithmus, der auf dem unüberwachten Lernen basiert, ist die Einteilung der aufgezeichneten Messpunkte in Arbeits- und Wendepunkte. Hierfür kann beispielsweise der k-means-Algorithmus genutzt werden, der im Abschnitt 2 beschrieben worden ist. Dies kann zur Optimierung von Arbeitseinsätzen oder auch für die Erstellung von Lastkartierungen sowie Zeitgliederungen genutzt werden.

In der Abbildung 3 wurde eine GPS-Spur dargestellt, bei der die Arbeits- und Wendepunkte mit Rauten und Kreisen markiert worden sind. Vorteil bei dieser Methode ist, dass auch das Vorgewende erkannt wird und man somit unabhängig von der Fahrtrichtung ist, wenn es um die Erkennung der eigentlichen Arbeit geht. Hierbei wurde dem Algorithmus vorgegeben, zwei Gruppen zu erzeugen und Daten anhand der Geschwindigkeit und des Kraftstoffverbrauchs zu den beiden Gruppen zuzuordnen. Zwar kann der Algorithmus auch ohne den Kraftstoffverbrauch die Arbeit einteilen, jedoch wird das Ergebnis damit deutlich robuster. Dieses Vorgehen kann unabhängig von der Maschinenart oder der Feldarbeit genutzt werden, liefert aber die besten Ergebnisse bei der Bodenbearbeitung, da dort die Spanne des Kraftstoffverbrauchs am größten ist.

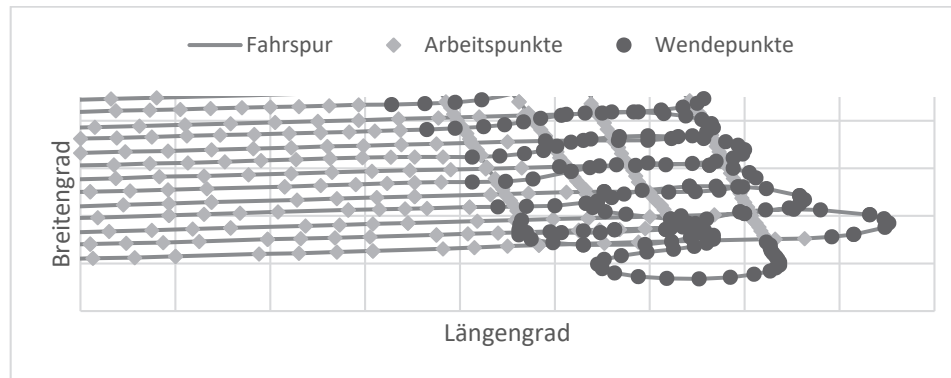


Abb. 3: Darstellung einer GPS-Spur, bei dem die Arbeits- und Wendepunkte mittels des k-means-Algorithmus in Gruppen eingeteilt worden sind

Da die Gruppen und deren Abgrenzung zueinander automatisch gebildet werden, muss keine zusätzliche Bedatung einer Funktion erfolgen, es wird direkt gruppiert. Zur Validierung kann die Position des Heckkrafthebers genutzt werden. Die Analyse von 150 Feld-einsätzen hat ergeben, dass der k-means-Algorithmus die Punkte mit einer Trefferquote von 85 % als Arbeits- oder Wendepunkte erkennt. Er hat jedoch eine hohe Standardabweichung, da die Trefferquote stark von der Qualität der Daten und der Art des Einsatzes abhängt.

5 Diskussion und Fazit

In diesem Beitrag wurde eine Auswahl von Algorithmen präsentiert, die auf dem unüberwachten Lernen basieren. Diese bringen verschiedene Vorteile, aber auch einige Nachteile mit sich. Die Algorithmen selbst sind meist weniger komplex als solche, die das überwachte Lernen nutzen, sie müssen nicht aufwendig mit Trainingsdaten initialisiert werden und haben meist auch einen geringeren Rechenaufwand. Sie können genutzt werden, um Anomalien in Datensätzen aufzuzeigen, jedoch muss der Anwender oftmals über Fachwissen zu den Daten verfügen, in dem die Algorithmen angewendet werden. Für eine erste schnelle Analyse der Daten können sie aber gut genutzt werden.

Einige der Algorithmen lassen sich auch automatisiert für die Auswertung nutzen wie etwa der DBSCAN, um Schläge zu erkennen. Die bisher gefundenen Schwächen liegen bei Schlägen, die nah beieinander liegen und somit das Risiko bergen, zusammengefasst zu werden. Auch bringt die Bedatung der Funktion und Validierung einen nicht unerheblichen Aufwand mit sich.

Der K-means-Algorithmus ist eine gute Option für eine schnelle und mit einem geringen Rechenaufwand verbundene Analyse der Feldarbeit. Die gezeigte Anwendung des Algorithmus bringt viele Vorteile für die Prozessanalyse, und die Einteilung in Arbeits- und Wendepunkte kann für verschiedene anschließende Auswertungen genutzt werden.

Das unüberwachte Lernen bringt aber auch Nachteile mit sich. Eine vollautomatische Anwendung der Algorithmen ist nur selten möglich und bedarf umfassender Testung, um robuste Ergebnisse zu liefern. Eine hohe Genauigkeit lässt sich nur schwer erreichen und bei vielen Anwendungen muss der Benutzer das Ergebnis anschließend noch begutachten. Somit ist oftmals eine Anwendung in der Forschung angebracht.

Es gibt noch weitere Algorithmen dieser Art. Diese werden im Rahmen von BiDa-LAP II sowie OskoNa auf ihre Einsatzmöglichkeiten analysiert und getestet werden.

Literaturverzeichnis

- [Br21] Brinkschulte, L.: Assistenzsysteme zur Reduktion des Schädigungsverhaltens von Komponenten einer mobilen Arbeitsmaschine. KIT Scientific Publishing, 2021.
- [Es96] Ester, M. et al.: A Density-Based Algorithm for Discovering Clusters in Large Spatial Databases with Noise: Proceedings / Second International Conference on Knowledge Discovery & Data Mining. AAAI Press, Menlo Park, Calif., S. 226-231, 1996.
- [He14] Heizinger, V. J.: Algorithmische Analyse von Prozessketten in der Agrarlogistik, Weihenstephan, 2014.
- [Jo10] Jolliffe, I. T.: Principal component analysis. Springer, New York, 2010.
- [Jo21] Jo, T.: Machine Learning Foundations. Supervised, Unsupervised, and Advanced Learning. Springer International Publishing; Imprint Springer, Cham, 2021.
- [Ma03] MacKay, D. J. C.: Information theory, inference and learning algorithms. Cambridge University Press, Cambridge, 2003.
- [Pi12] Pierre Baldi: Autoencoders, unsupervised learning, and deep architectures. In (ICML Hrsg.): Workshop on Unsupervised and Transfer Learning, 2012.
- [PEN20] Poteko, J.; Eder, D.; Noack, P. O.: Bestimmung des Betriebsmodus landwirtschaftlicher Maschinen auf Basis von GNSS-Messwerten. In (Leibniz-Institut für Agrartechnik und Bioökonomie e.V Hrsg.): Lecture Notes in Informatics (LNI), Bonn, S. 241-246, 2020.
- [SM18] Stein, T.; Meyer, H. J.: Automatic machine and implement identification of an agricultural process using machine learning to optimize farm management information systems. In (Leibniz-Institut für Agrartechnik und Bioökonomie e.V Hrsg.): Bornimer Agrartechnische Berichte 101, Potsdam, S. 19-26, 2018.