

Multidimensionale Datenmodellierung und Analyse zur Qualitätssicherung in der Fertigungsautomatisierung

M.Sc. **B. Lindemann**, Dr.-Ing. **N. Jazdi**, Prof. Dr.-Ing. **M. Weyrich**,
Institut für Automatisierungstechnik und Softwaresysteme, Universität
Stuttgart, Stuttgart;

Kurzfassung

Unerwartete Ausfälle von Anlagenkomponenten sowie unvorhersehbare Prozessereignisse und Anomalien sind Treiber erhöhter Ineffizienzen in Form von Stillstandzeiten und schwankender Produktqualität. Produzierende Unternehmen stehen vor der Herausforderung, trotz entlang der gesamten Wertschöpfungskette auftretender Dissonanzen, ein hochwertiges und reproduzierbares Qualitätsergebnis zu gewährleisten. Dazu werden Lösungen benötigt, die die zunehmende Komplexität entlang vernetzter und hochgradig dynamischer Prozessketten beherrschbar machen. Intelligente Sensor-Netzwerke erlauben den Aufbau umfangreicher Infrastrukturen zur Datenerfassung und sind Wegbereiter für den Einzug datenanalytischer Konzepte und Lösungen in die Fertigungsautomatisierung ('Big Data Analytics').

In diesem Beitrag wird ein datengetriebener Ansatz zur Qualitätsoptimierung vorgestellt, auf dessen Basis Interdependenzen entlang von Prozessketten erfasst und mit Hilfe von multidimensionalen Datenmodellen abgebildet werden können. Aufgrund der Tatsache, dass aktuelle 'Self-Learning'-Ansätze potenzielle Interdependenzen zwischen Prozessen nicht adäquat modellieren, mangelt es diesen aufgrund der isolierten Optimierung von Einzelprozessen an Skalierbarkeit. Systeme und Prozessketten mit hochdimensionalen Datenräumen können dadurch nicht oder nur unzureichend abgebildet werden. Durch eine verteilte, multidimensionale Modellierung ist dagegen die Integration multipler Prozessschritte in die analytische Betrachtung möglich. Das in dieser Arbeit beschriebene Konzept stellt darüber hinaus einen „Grey-Box-Ansatz“ dar, indem es vorsieht, bestehendes Wissen aus dem Engineering um aus Daten extrahierte Modellstrukturen zu erweitern und zu einem ganzheitlichen Modell zusammenzuführen.

Darauf aufbauend wird das in der vorliegenden Arbeit propagierte Gesamtsystem durch zwei zusammenwirkende Rückkopplungsschleifen realisiert. Bestehendes Wissen wird in einer direkt steuernd auf die Prozesskette einwirkenden Schleife mit Hilfe von Regelstrukturen abgebildet, die die verteilten, multidimensionalen Datenmodelle semantisch verknüpfen.

Diese wird um eine vollständig datengetriebene Schleife ergänzt, die die Dimension des erfassten Datenraumes reduziert, Merkmale auf Basis maschinellen Lernens extrahiert und die Wissensbasis kontinuierlich erweitert. Dadurch wird eine ganzheitliche und datenbasierte Qualitätsoptimierung realisiert. Der Ansatz wird momentan prototypisch für eine Modell-Prozesskette der Massivumformung unter Laborbedingungen sowie für zwei Pilotanlagen im realen industriellen Umfeld der Unternehmen Otto Fuchs KG und Hirschvogel Automotive Group umgesetzt. In diesem Beitrag wird die Datenextraktion und -aufbereitung, die prototypisch umgesetzte Datenmodellierung, der Ansatz zur Datenanalyse sowie das Gesamtkonzept vorgestellt.

1. Einleitung

Für die industrielle Qualitätssicherung ergeben sich durch die Digitalisierung und die Potentiale lernfähiger, intelligenter Systeme zunehmend neue Möglichkeiten. Dieser Beitrag stellt einen Ansatz vor, der sich von bisherigen Ansätzen für selbst-lernende Systeme dadurch unterscheidet, dass eine verteilte, multidimensionale Modellierung für die Integration multipler Prozessschritte in die analytische Betrachtung genutzt wird. Die nachfolgende Arbeit ist wie folgt gegliedert: Abschnitt 2 gibt einen Überblick über den Stand der Technik für wissensbasierte und selbst-lernende Systeme und stellt das Gesamtkonzept vor. Abschnitt 3 führt in das Konzept zur Datenerfassung und –verarbeitung ein. Der verwendete Ansatz zur Datenmodellierung und die vorgeschaltete Datenintegration werden in Abschnitt 4 vorgestellt. In Abschnitt 5 wird das maschinelle Lernverfahren erläutert, dass zur Detektion von Anomalien und für die Lernfähigkeit der Qualitätssicherung verwendet wird. Abschnitt 6 schließt die vorliegende Abhandlung mit einer Darstellung der Realisierung des beschriebenen Assistenzsystems.

2. Stand der Technik

Wissensbasierte Systeme sind in der Fertigungsautomatisierung und Qualitätssicherung weit verbreitet. Diese Systeme werden häufig in Form von Inferenzsystemen realisiert, da Expertenwissen und Wissen aus dem Engineering integriert werden kann. In den letzten Jahren sind verschiedene anwendungsorientierte Ansätze und Erweiterungen derselben entstanden. [1] hat Verfahren des maschinellen Lernens mit temporären Fuzzy-Inferenz-Systemen kombiniert und dadurch einen neuen Ansatz für die Qualitätskontrolle und Fehlerdiagnose von Fertigungssystemen geschaffen. Um das bestehende Wissen zu erweitern wird darauf abgezielt, neues Wissen aus den großen Mengen von Prozessdaten zu gewinnen, die heutzutage in vielen produzierenden Betrieben aufgenommen werden. Der

in [2] beschriebene Ansatz für ein Assistenzsystem zur Überwachung und Diagnose automatisierter Systeme propagiert zwei Anwendungsschichten, eine wissensbasierte Schicht und eine Schicht des maschinellen Lernens. Eine mit Hilfe von Experten erarbeitete Regelbasis wird um die Erkenntnisse einer kontinuierlich mit Prozessdaten gefütterten Bayes'schen Klassifikation verfeinert. [3] präsentiert ein sich selbst optimierendes System für industrielle Abtragungsprozesse. Die Selbstoptimierung wird dabei auf Basis einer Support-Vector-Machine umgesetzt. In [4] wird ein auf Neuronalen Netzstrukturen basierender Ansatz vorgestellt, wobei die Netze in ein Inferenzsystem zur Diagnose von Transformatoren integriert werden. Die Wissensbasis des Systems wird zyklisch basierend auf der Dichtestruktur der aufgezeichneten Daten aktualisiert. Die Literatur zeigt verschiedene Ansätze für Assistenzsysteme, die Expertenwissen mit rein datengetriebenen Konzepten kombinieren, die jedoch in der Mehrzahl der Fälle für einzelne Prozesse oder Automatisierungssysteme ausgelegt sind, z.B. [5].

Der Hauptbeitrag der vorliegenden Arbeit besteht in der Konzeption und Umsetzung eines Assistenzsystems zur Qualitätssicherung, das im Gegensatz zu bestehenden 'Self-Learning'-Ansätzen potenzielle Interdependenzen zwischen einzelnen Automatisierungssystemen und Prozessen modelliert. Dabei gibt es diverse Herausforderungen, die es zu bewältigen gilt. Für ein prozessübergreifendes Eingreifen entlang industrieller Prozessketten müssen Daten von einzelnen Messeinheiten abgerufen und auf einem übergeordneten Server, auf dem das Assistenzsystem realisiert ist, zusammengeführt werden.

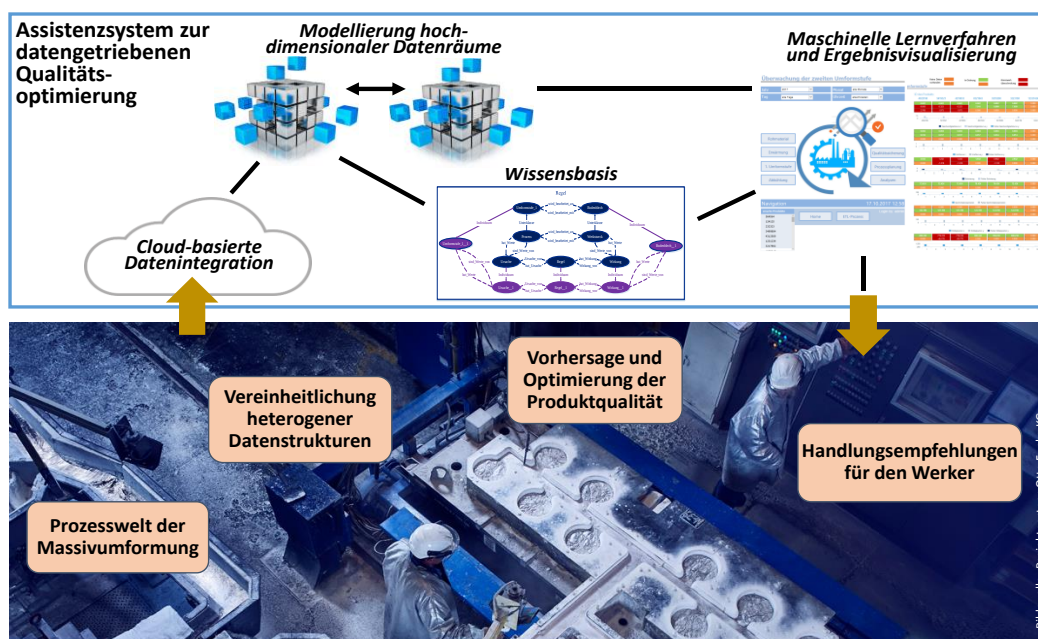


Bild 1: Assistenzkonzept und Systemarchitektur

Daher wird ein Datenerfassungskonzept benötigt, das der Heterogenität der Datenquellen in der Feldebene Rechnung trägt. Auf dieser Basis wird in der vorliegenden Arbeit ein Cloud-basierter Ansatz verfolgt, der ein multidimensionales Modellierungskonzept zur Verarbeitung der Daten beinhaltet. Engineering-Wissen wird in regelbasierten Metamodellen formalisiert, die zur Inferenz genutzt werden. Dieses Wissen über die Prozesse und ihre Abhängigkeiten entlang der Prozesskette wird kontinuierlich über Lernprozesse verfeinert. In dieser Arbeit wird der k-Means-Algorithmus genutzt, um das Auftreten von Prozessanomalien zu lernen und die Prädiktionsgüte derselben zu optimieren. Das Gesamtkonzept ist in Abbildung 1 dargestellt. Es wird anhand von Prozessketten der Massivumformung realisiert und evaluiert.

3. Steuerungs-basierte Datenerfassung und -verarbeitung

Für eine datengetriebene Optimierung der Produkt- und Prozessqualität ist die sichere Zuordnung von Prozess- und Bauteileigenschaften essentiell. Bei der Rückverfolgung einzelner Werkstücke oder Chargen entlang von Prozessketten stellt die Erfassung und Zuordnung von Prozessdaten zu den entsprechenden Losen eine große Herausforderung in der industriellen Umsetzung dar. Darüber hinaus stellt die Heterogenität der Datenquellen in Bezug auf Datenformate, Protokolle und Abstraten eine weitere Problematik dar. Zur Lösung der aufgezeigten Probleme wurde ein Softwarekonzept entwickelt, das die gesammelten Daten einheitlich in die Cloud überträgt. Das Konzept ist angelehnt an das von [6] publizierte Vorgehen. Das entwickelte Modul besteht aus einer individuellen Schnittstelle für jede Datenquelle in der Feldebene und einer standardisierten Schnittstelle zur Weiterleitung der aufgezeichneten Daten. Die einzelnen, individuell auf die jeweilige Datenquelle angepassten Schnittstellenmodule laufen auf Steuergeräten, die eine einheitliche, in der Cloud realisierte Schnittstelle ansprechen. Die für die Steuergeräte entwickelte Datenerfassungskomponente ist als Zustandsautomat realisiert, der durch seinen generischen Aufbau flexibel in bestehenden SPS-Code eingebunden werden kann. Die zyklisch am Bussystem und der entsprechend anliegenden Peripherie abgefragten Sensordaten werden auf ein einheitliches REST-Modell abgebildet. Aufgrund der Tatsache, dass der größte Teil der Prozessdaten, die an Fertigungssystemen aufgezeichnet werden, Zeitreihendaten sind, besteht das REST-Modell grundsätzlich aus den Komponenten Zeitinformation, Metadaten und Daten. Die gemessenen Prozessparameter werden in Abhängigkeit der Protokollstruktur konvertiert, auf das Modell abgebildet und an die standardisierte Schnittstelle weitergegeben. In der Cloud findet die weitere Verarbeitung der Daten in einer Echtzeitdatenbank statt, die es ermöglicht Ad-hoc-Berechnungen durchzuführen und Analyseverfahren gemäß dem Online Analytical Processing (OLAP) zu

implementieren. Diese Systemkomponente wird für die eindeutige Zuordnung von Werkstücken und Prozessparametern genutzt und ermöglicht je nach produktionstechnischen Umwelt- und Rahmenbedingungen eine Online-Verfolgung von Einzelteilen bzw. Chargen.

4. Multidimensionale Datenmodellierung und Bauteilrückverfolgung

Um eine effiziente Qualitätssicherung realisieren zu können, werden OLAP-Modelle auf Grundlage der Datenräume der einzelnen Prozessschritte erstellt [7]. Dazu wurde ein ETL-Stack (Extract-Transform-Load) implementiert, der die Daten der verschiedenen Datenquellen extrahiert und durch Transformationen in ein strukturiertes Modell überführt, das In-Memory gehalten und konsolidiert wird. Die automatisierte Datenaufbereitung und -verarbeitung ist in Abbildung 2 dargestellt.

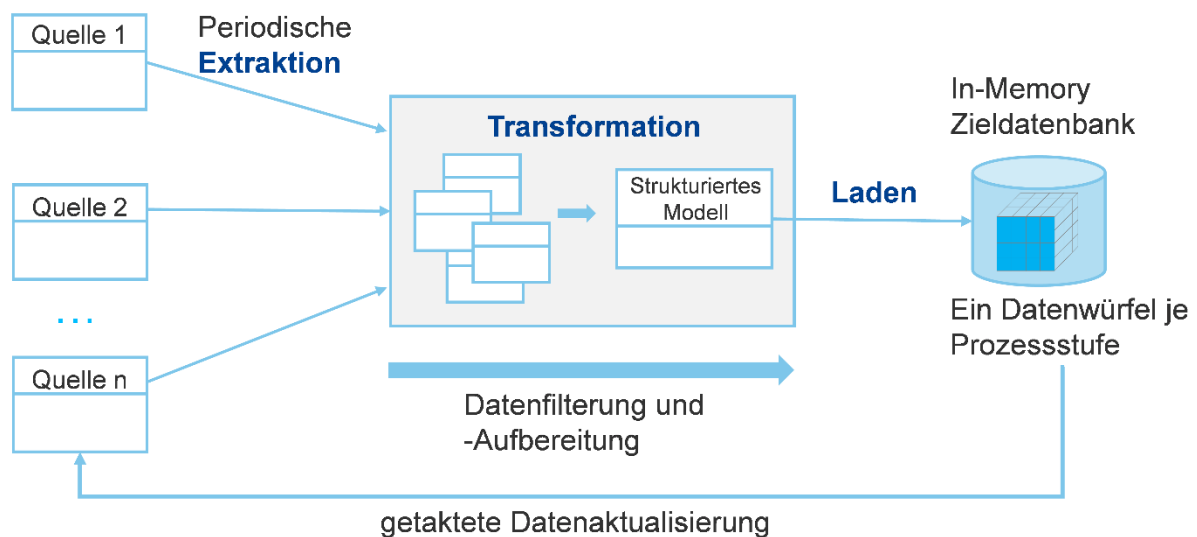


Bild 2: ETL-basiertes Konzept zur Homogenisierung des Datenraumes

Die Modelle werden aus einer Vielzahl an Dimensionen aufgebaut. Dimensionen können bspw. durch die Zeit-, Orts- oder Produktinformationen aufgespannt werden. Diese enthalten eine interne hierarchische Struktur, in der jede Verzweigung durch ein Element beschrieben wird. Basis-Elemente sind Vektoren von Prozessdaten zugeordnet und markieren die innerste Schicht des Datenmodells. Ein Basiselement könnte z.B. die ID 1234 eines Werkstücks sein. Höhere Schichten kapseln kumulierte Informationen der unteren Schichten. Eine höhere Hierarchieebene der Dimension ‚Produktinformation‘ würde bspw. die Chargeninformation beinhalten. Der ETL-Stack bildet neue Daten auf dieses hierarchisch strukturierte Modell ab. Wird bspw. ein neues Produkt eines neuen Produkttypen produziert

oder Chargen anstelle von einzelnen Bauteilen verfolgt, können diese Informationen auf die Dimensionsstruktur projiziert und nach einem Abgleich hinzugefügt werden. Diese Projektion bewirkt die Erstellung von Elementen innerhalb der Hierarchie wie zum Beispiel die Erstellung höherer Schichten für die Chargenverfolgung. Folglich können unterschiedliche Prozessdatenstrukturen unterschiedlicher Produkte und Mengen angepasst werden. Die gleiche Flexibilität gilt für jedes andere Element jeder anderen Ebene in der hierarchischen Dimensionsstruktur. Die zugehörigen Daten werden durch die Anwendung geeigneter Aggregationsverfahren berechnet. Somit können heterogene Daten in ein homogenes Modell integriert werden, da es möglich ist, verschiedene Schichten mit unterschiedlichen Granularitäten zu assoziieren. Die Rückverfolgung und Überwachung einer beliebigen Losgröße kann dadurch bereits durch das einfache Navigieren durch verschiedene Ebenen des hierarchischen Metadatenkonstrukts erreicht werden [8]. Der Modellierungsansatz ist in Abbildung 3 als Würfel visualisiert.

N-dimensionale Modellstrukturen zeichnen sich dadurch aus, dass sie beliebig zusammengesetzt und reduziert werden können [9]. So können die Datenmodelle einzelner Prozessschritte zum Gesamtmodell der Prozesskette komponiert werden. Darüber hinaus werden zu Modellen Metamodelle erzeugt, die das Expertenwissen kapseln indem sie Schwellwerte und Regeln für Prozessparameter beschreiben und in einem formalen Modell abbilden. So kann eine effektive Prozessüberwachung und –diagnose umgesetzt werden.

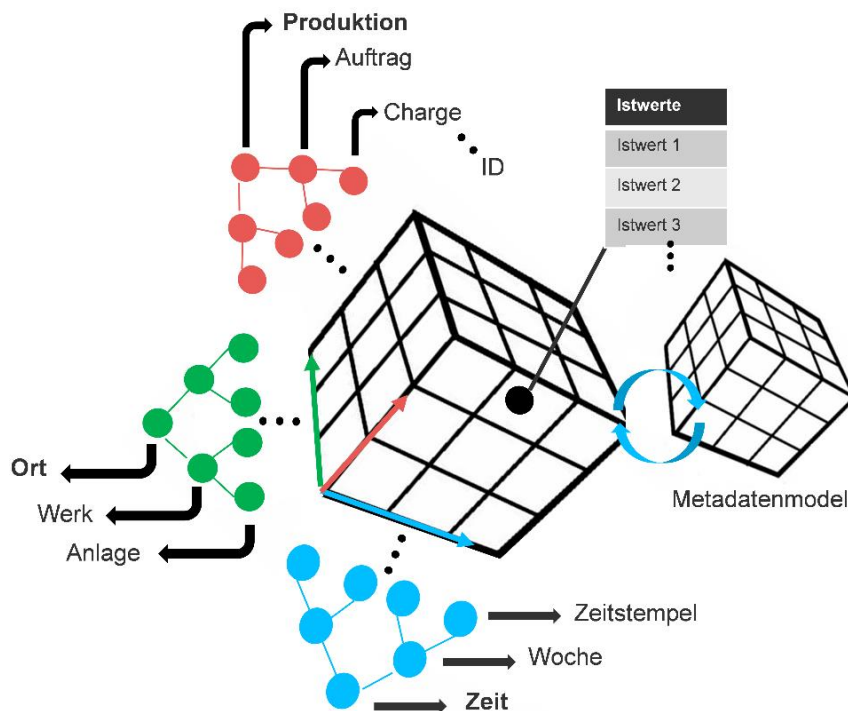


Bild 3: Multidimensionale Modellierung von Prozessdaten und Bauteileigenschaften

Im Folgenden wird ein Beispiel für die Darstellung von Expertenwissen gegeben: Es sei angenommen, dass die Umgebungsbedingungen entlang der betrachteten Prozesskette zu einer Veränderung hinsichtlich der Qualität des verwendeten Kühlmittels führen, das wiederum die Qualität von Prozess und Produkt beeinflusst. Die Reduktion in der Qualität der Kühlmittel hängt von zahlreichen Faktoren ab, die durch den Experten beschrieben werden können. Dazu gehören unter anderem die Verwendungsdauer (Zeitdimension) sowie der Maschinentyp (Ortsdimension), auf der das Kühlmittel eingesetzt wurde. Folglich können diese Nebenwirkungen formalisiert, auf die Dimensionsstruktur abgebildet und in Form eines Metamodells umgesetzt werden. Eine automatisierte Prozessüberwachung auf Basis von Expertenwissen ist damit möglich.

5. Maschinelles Lernen zur Detektion von Anomalien

Um das beschriebene System um eine intelligente Rückkopplungsschleife zum Prozess zu erweitern, wurden verschiedene algorithmische Ansätze untersucht, die eine frühzeitige Detektion von Prozessanomalien erlauben und damit zu einer verbesserten Produktqualität beitragen. Nachfolgend wird ein auf dem k-Means Clustering Algorithmus basierendes Verfahren vorgestellt. Der k-Means Algorithmus stellt ein unüberwachtes Lernverfahren dar. Daten mit ähnlichen Eigenschaften werden durch das Verfahren in Clustern zusammengefasst. Die Anzahl der möglichen Cluster, muss vor dem Training definiert werden. Jedes Datum ist in den meisten Fällen durch eine Vielzahl an Merkmalen charakterisiert. Dadurch entsteht ein hochdimensionaler Eingangsdatenraum in dem jeder Datenpunkt durch einen Vektor dargestellt werden kann. Der Algorithmus reduziert den euklidischen Abstand zwischen den Clusterzentren $\mu_1, \mu_2, \dots, \mu_k$ und den erfassten Prozessdaten $x^{(1)}, x^{(2)}, \dots, x^{(m)}$, sodass insgesamt folgendes Optimierungsproblem zu lösen ist: $c^{(i)} := \arg \min_j \|x^{(i)} - \mu_j\|^2$. Der Index k entspricht der Anzahl der Cluster und der Index m der Anzahl der erfassten Datensätze [10]. In der vorliegenden Studie wurde mit einer zeitsensitiven Variante des k-Means Algorithmus gearbeitet, indem der Merkmalsraum auf Basis einer Sliding-Window-Funktion künstlich vergrößert wurde. Dadurch kann die Anzahl der Trainingsmerkmale bis zu einem sinnvollen Betrachtungszeitraum vergrößert werden. Die Vergrößerung der Batchgröße kann jedoch nicht beliebig erfolgen, weil dies zu diversen Nachteilen wie einer erhöhten Trainingszeit führt [11].

Die Anzahl der potentiellen Cluster ist selten bekannt, wird aber als Eingangsparameter zur Durchführung des Trainings auf den gemessenen Sensordaten benötigt. Eine heuristische

Methodik, mit der die Anzahl der Cluster bestimmt werden kann, bietet die Ellenbogenmethode. Dabei wird der k-Means Algorithmus wiederholt mit demselben Datensatz ausgeführt, wobei die Clusteranzahl variiert wird. Der mittlere euklidische Abstand zu den Clusterzentren wird als Homogenitätsmetrik genutzt. Die optimale Anzahl der Cluster wird durch eine grafische Auswertung bestimmt. Das Training wird auf einem großen, repräsentativen Datensatz durchgeführt, der das Normalverhalten des zu Grunde liegenden Prozesses widerspiegelt. Die eingelernte Clusterstruktur kann damit als Gut-Modell aufgefasst werden. Anschließend wird das trainierte Modell verwendet, um den Abstand zwischen einem Testdatensatz, bestehend aus neu aufgenommenen Daten, und den Clusterzentren zu berechnen. Anomalien im Testdatensatz lassen sich über eine definierte Grenzwertüberschreitung an Testdatenpunkten mit großer Entfernung zu den Clusterzentren sowie über die Entstehung neuer, zuvor nicht vorhandener Cluster detektieren. Die Funktionalität der datengetriebenen Systemkomponente ist in Abbildung 4 in Form eines Sequenzdiagramms veranschaulicht. Die aus der Datenanalyse gewonnen Erkenntnisse werden darüber hinaus in die bestehende Wissensbasis integriert. Detektierte Anomalien werden dem Anlagenbetreiber auf der grafischen Benutzeroberfläche des Assistenzsystems visualisiert.

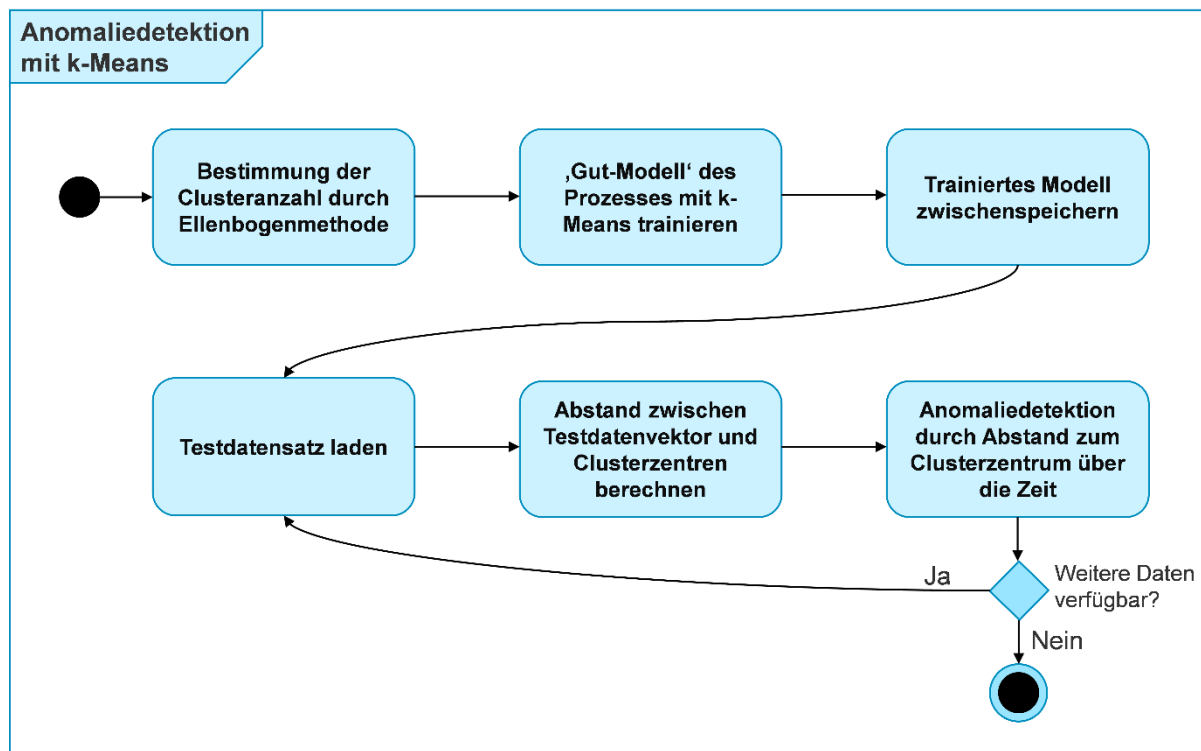


Bild 4: Sequenzdiagramm zur Anwendung des k-Means-Algorithmus

6. Realisierung und Evaluation der beschriebenen Funktionalitäten

Um die Systemarchitektur zu testen, werden sowohl die Simulation einer Prozesskette der Massivumformung sowie eine reale Prozesskette verwendet. Eine Beschreibung der Prozesskette und seiner umformtechnischen Besonderheiten kann in [12] nachgelesen werden. Die Prozessdaten werden in Zyklen von einer Millisekunde aus den Steuerungen extrahiert, um den Verlust relevanter Prozessinformationen aufgrund einer zu niedrigen Abtastrate zu verhindern. Der Datenraum umfasst 18 verschiedene Sensorwerte von 5 aufeinander folgenden Prozessschritten. Die standardisierte Schnittstelle läuft in der Cloud und leitet die Daten im JSON-Format an die OLAP-Datenbank weiter. Die Dateigröße ist ein Eingabeparameter des Zustandsautomaten und kann vom Benutzer zur Vermeidung eines Buffer Overflows angepasst werden. Der gesamte Datenintegrationsprozess wird zyklisch und GPU-unterstützt ausgeführt, um paralleles Rechnen zu ermöglichen. Dadurch kann eine effiziente Online-Überwachung realisiert werden. Detaillierte Informationen zur verwendeten Hardware sind in Tabelle 1 zu finden.

Tabelle 1: Verwendete Hardware

OS	GPU	CPU	RAM	Netzwerk
CentOS Server Version 7.3	2x Tesla K80 24 GB GDDR5	18-Core @ Xeon 2.10 GHz	512 GB DDR4 LRDIMM	1 Gigabit Ethernet

Der Einsatz einer leistungsfähigen Grafikkarte ermöglicht neben den Vorteilen für die Datenmodellierung auch eine erhebliche Verkürzung der Trainingsdauer des Cluster-Algorithmus. Das liegt daran, dass sich der Algorithmus aus einer Vielzahl einfacher und größtenteils unabhängiger Rechenoperationen zusammensetzt, sodass mit Hilfe einer GPU eine effiziente Parallelisierung der Rechenoperationen umgesetzt werden kann. Zu Beginn der Trainingszeit werden die vorverarbeiteten Datensätze in den Arbeitsspeicher geladen. Durch eine Vergrößerung der Merkmalsanzahl konnte der Geschwindigkeitsvorteil weiter erhöht werden. Im Hinblick auf die betrachtete Prozesskette kann die Anzahl der Merkmale jedoch nicht beliebig vergrößert werden, da die Größe des Sliding-Window von dem zu analysierenden Szenario abhängig ist. Darüber hinaus wurde bei einigen Operationen ein starker Performanceeinbruch beobachtet, der auf eine erhöhte Anzahl an Speicherzugriffen zurückzuführen ist. Abbildung 5 zeigt das Ergebnis der Detektion einer Anomalie bei Pumpe 5 einer Räderpresse. Pumpe 4 und 8 erfüllen Spezialaufgaben und werden daher nicht als Referenz herangezogen. Abbildung 6 stellt die Benutzeroberfläche des Assistenzsystems dar. Es ist als Web-Interface realisiert, sodass ein flexibler Zugriff von mobilen Endgeräten ermöglicht wird.

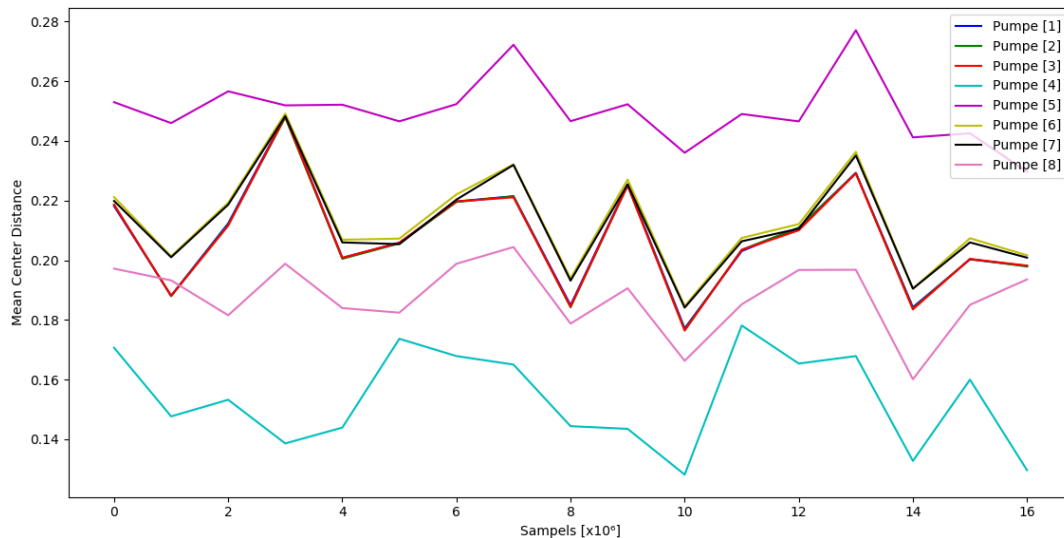


Bild 5: Qualitätssicherung durch Detektion von Anomalien

Die links dargestellte Ansicht zeigt eine anmeldebasierte Überwachung der erfassten Prozessparameter. Die Färbung kommt durch die Anwendung der aus Expertenwissen gewonnen Regeln zustande, wobei eine kritische Überschreitung eines Grenzwertes durch ein rotes aufleuchten signalisiert wird. Durch die zu Grunde liegende Datenmodellierung kann der Benutzer in der Ansicht in den drei Hauptdimensionen Zeit, Ort und Produkt navigieren und ist dadurch in der Lage einzelne Werkstücke oder Chargen zu einem bestimmten Zeitpunkt an einem bestimmten Prozess zu überwachen und die entsprechende Qualität zu gewährleisten. Die zweite in der untenstehenden Abbildung dargestellte Ansicht ermöglicht den Vergleich aufeinanderfolgend hergestellter Werkstücke in eines jeden Prozessschrittes. Abweichungen werden durch den stetigen Abgleich der Daten mit dem eingelernten Cluster-Modell detektiert und angezeigt. Trajektorien und schleichende Abweichungen können erkannt und eine effiziente Qualitätsüberwachung der Prozesse und Produkte umgesetzt werden.

7. Zusammenfassung und Ausblick

In dieser Arbeit wurde ein Ansatz für ein Assistenzsystem vorgestellt, der es ermöglicht, Prozess- und Produktdaten über eine Cloud-basierte Datenintegration zusammenzuführen. Die extrahierten Daten werden über einen ETL-Stack in ein homogenes und strukturiertes Modell überführt. Das Konzept der multidimensionalen Datenmodellierung ermöglicht die flexible Navigation in den Daten. Die Integration von Expertenwissen in Form von anwendungsspezifischen Regeln ermöglicht Assistenzfunktionen wie die Diagnose Überwachung von Prozessereignissen.



Bild 6: Benutzeroberfläche des Assistenzsystems

Die beschriebene wissensbasierte Komponente wird um rein datengetriebene Verfahren erweitert und befähigt das System zu weiteren Assistenzfunktionen wie der Vorhersage oder der frühzeitigen Detektion von Anomalien. Das Konzept wurde anhand einer am Institut für Umformtechnik der Universität Stuttgart bestehenden Prozesskette entwickelt und anhand realer und simulierter Komponenten getestet. Im weiteren Verlauf wird die Systemarchitektur an einer vollständig automatisierten Produktionslinie getestet bevor das Konzept auf zwei Pilotanlagen der industriellen Partner Otto Fuchs KG sowie Hirschvogel Automotive Group übertragen wird. Die Pilotanlagen dienen dazu, die Übertragbarkeit und Anwendbarkeit des informationstechnischen Ansatzes im realen industriellen Umfeld zu zeigen. Das Ziel besteht darin, die Möglichkeiten und Potentiale zur Erhöhung der Gesamtanlageneffektivität entlang der gesamten Wertschöpfungskette am Beispiel der Massivumformung aufzuzeigen.

8. Literaturangaben

- [1] Mahdaoui, R., Mouss, L.H., Kadri, O. and et al. (2012). Fault prognosis by temporal neuro-fuzzy systems: application for manufacturing systems. In Proceedings of IEEE SETIT, Sousse, Tunisia, 2012, pp. 1–6.
- [2] Abele, L., Anic, M., Gutmann, T., Folmer, J., Kleinstaub, M. and Vogel-Heuser, B. (2013). Combining knowledge modeling and machine learning for alarm root cause analysis. In 7th IFAC MIM, St. Petersburg, Russia, June 19-21, 2013.

- [3] Denkena, B. and Dittrich, M. (2016). Self-Optimizing cutting process using learning process models. In 3rd Conference on System-Integrated Intelligence (SysInt), Procedia Technology 26 (2016) pp. 221 – 226.
- [4] Fischer, M. and Tenbohlen, S. (2009). Improved condition assessment by fuzzy modelling, adjustment and merging of DGA's interpretation methods. In Conference Proceedings of ISH, August 24-28, Cape Town, South Africa, no. F-16, pp. 1556-1562.
- [5] Endelt, B. and Volk, W. (2013). Designing and iterative learning control algorithm based on process history – using limited post process geometrical information. In IDDRG 2013 Conference Proceedings, Zurich, Switzerland, June 2-5, pp. 69-74.
- [6] Faul, A., Jazdi, N., Weyrich, M.: Approach to interconnect existing industrial automation systems with the industrial internet. In 21st IEEE International Conference on Emerging Technologies and Factory Automation (ETFA) 2016.
- [7] Schwarz, H.: Konzeptueller und logischer Data-Warehouse-Entwurf: Datenmodelle und Schematypen für Data Mining und OLAP. In: Informatik Forschung und Entwicklung. 18(2004) 2, S. 53-67.
- [8] Lindemann, B., Jazdi, N., Weyrich, M.: Softwaresysteme zur Qualitätssicherung in der Umformtechnik - Ein Ansatz für die echtzeitfähige und prozessübergreifende Qualitätsüberwachung. In Industrie 4.0 Management 33 (2017) 6.
- [9] Shin, S.J., Woo, J., Rachuri, S.: Predictive analysis model for power consumption in manufacturing. In 21st CIRP Conference on Life Cycle Engineering, Procedia CIRP 15, 2014.
- [10] Piech, C.: K Means. (2013), URL: <http://stanford.edu/~cpiech/cs221/handouts/kmeans.html> (Stand: 20.02.18).
- [11] Keskar, N. S., Mudigere, D., Nocedal, J., Smelyanskiy, M., Tang, P.T. P.: On Large-Batch Training for Deep Learning: Generalization Gap and Sharp Minima. In: ICLR 2017. 2017, S. 1-16.
- [12] Liewald, M.; Karadogan, C.; Felde, A.; Lodwig, R.: Entwicklung und Integration digitaler Technologien in Prozessfolgen der Massivumformung. In: Neuere Entwicklungen in der Massivumformung (NEMU) 2017.