



Mathematik revolutioniert die Supply Chain –
Analytics schafft neue Werte und Wettbewerbsvorteile

White Paper: Data Analytics in der Supply Chain

Andreas Bärmann,
Julius Mehringer,
Christian Menden,
Ursula Neumann,
Julia Schemm,
Oskar Schneider,
Benedikt Sonnleitner,
Markus Weissenbäck

Herausgeber:

Alexander Pflaum,
Roland Fischer

Download hier:



Fraunhofer-Arbeitsgruppe für
Supply Chain Services des Fraunhofer IIS



Herausgeber: Alexander Pflaum | Roland Fischer
Fraunhofer-Arbeitsgruppe für Supply Chain Services SCS des
Fraunhofer IIS

Data Analytics in der Supply Chain

Mathematik revolutioniert die Supply Chain – Analytics schafft neue Werte und Wettbewerbsvorteile

Autoren (alphabetisch):

Andreas Bärmann
Julius Mehringer
Christian Menden
Ursula Neumann
Julia Schemm
Oskar Schneider
Benedikt Sonnleitner
Markus Weissenböck

Inhalt

ABSTRACT	5
1 ANALYTICS ALS BEFÄHIGER IN EINER DIGITALISIERTEN SUPPLY CHAIN	7
2 ANWENDUNGSNAHE METHODEN DER GRUPPEN „DATA SCIENCE“ UND „OPTIMIZATION“ DER FRAUNHOFER ARBEITSGRUPPE FÜR SUPPLY CHAIN SERVICES	11
2.1 Data Science.....	11
2.1.1 Überwachtes Lernen.....	13
2.1.2 Unüberwachtes Lernen	18
2.2 Mathematische Optimierung	21
2.2.1 Heuristiken	22
2.2.2 Exakte Verfahren.....	23
3 ERFOLGSGESCHICHTEN AUS DER PRAXIS	27
3.1 Abgeschlossene Projekte	33
3.1.1 Fehlerquellen erkennen – Fehler vermeiden	33
3.1.2 Langfristig zuverlässig: Ersatzteilbedarfe für das nächste Jahrzehnt vorhersagen	35
3.1.3 Dynamische Lagerplanung bei Schnellecke	36
3.2 Laufende Projekte	27
3.2.1 KITE: Künstliche Intelligenz im Transport zur Emissionsreduktion.....	27
3.2.2 PRODAB: Durchlaufzeitprognose mit Bayes’schen Netzen	28
3.2.3 Ressourcenschonende und effiziente Produktionsplanung beim Teezulieferer Martin Bauer	32
4 WEITERE / KOMMENDE EXPERTISE IM BEREICH SUPPLY CHAIN ANALYTICS.....	39
4.1 Integration von Domänenwissen in Forecasting.....	39
4.2 Quantifizierung der Unsicherheit in Forecasting.....	40
4.3 Kausalitätsanalysen	42
4.4 Online-Optimierung und lernende Optimierungsverfahren	43
5 ANHANG.....	44
5.1 Literatur.....	44
5.2 Abbildungsverzeichnis	44
5.3 Tabellenverzeichnis	44
5.4 Die Autoren.....	44
6 LITERATURVERZEICHNIS.....	42

ABSTRACT

Im Supply Chain Management erfährt das Forschungsgebiet Data Analytics eine wachsende Aufmerksamkeit. In den letzten Jahren erfassen die Unternehmen in diesen Lieferketten immer durchgängiger die relevanten Prozessdaten und der verfügbaren Rechenleistung sind kaum noch Grenzen gesetzt. Eine Vision wird also langsam Wirklichkeit: Immer mehr Entscheidungen entlang der Lieferkette können (teil-) automatisiert auf Basis von Daten solide getroffen werden. Dabei helfen Analytics-Methoden. Dieses Whitepaper gibt eine Übersicht über gängige Analytics-Methoden und betrachtet ihre Anwendungsfälle im Supply Chain Management. In Beispielen aus der Praxis wird der Stand der Technik erläutert und ein Ausblick auf künftige Forschungs- und Entwicklungsfelder entlang der Wertschöpfungsketten gegeben. Bei den Methoden und ihre Anwendungen im Supply Chain Management wird unter Machine Learning und mathematischer Optimierung unterschieden.

1 ANALYTICS ALS BEFÄHIGER IN EINER DIGITALISIERTEN SUPPLY CHAIN

Die Definition der Supply Chain (Versorgungskette) umfasst die Planung, Durchführung und Kontrolle aller Aktivitäten im Zusammenhang mit dem Material- und Informationsfluss vom Einkauf der Rohstoffe über die Produktion bis zur endgültigen Lieferung des Produkts an den Kunden (Schulte 2009). Abgesehen von den physischen und monetären Aktivitäten in einer Supply Chain sind ihre Akteure hauptsächlich mit dem Treffen von Entscheidungen beschäftigt. Üblicherweise ist das Ziel der Entscheidungsträger, die physischen Prozesse so effektiv und effizient wie möglich zu gestalten. Die getroffenen Entscheidungen haben dabei unterschiedliche individuelle Tragweiten und sind mit schwankenden Unsicherheiten behaftet. Beispielsweise entscheiden Disponenten über die Menge an Gütern, die beschafft wird, der Vertrieb entscheidet über den Preis, den ein Kunde zahlen soll, und das Management entscheidet über die Ausrichtung ganzer Geschäftsbereiche.

Um beim Treffen dieser Entscheidungen zu unterstützen oder diese sogar zu automatisieren, werden vermehrt auch digitale Lösungen eingesetzt. Die Hoffnung, mithilfe der Digitalisierung mathematisch komplexe Entscheidungen künftig datengetrieben treffen zu können, erklärt vielleicht die Strahlkraft von Begriffen wie „Künstliche Intelligenz“, „Big Data“ oder „Deep Learning“. Alle diese Ansätze versuchen die Welt (oder spezifische Unternehmensprozesse) mathematisch zu beschreiben, um dann die beste Lösung für eine gegebene Problemstellung auszurechnen. Die Anwendungen von Business-Intelligence-Lösungen der 2000er und 2010er Jahre hatten den Anspruch Informationen aufbereitet darzustellen. Der nächste Schritt in Richtung Digitalisierung ist es nun Software zu entwickeln, welche Vorschläge für Lösungen generiert, die den Menschen dabei unterstützt fundierte Entscheidungen zu treffen. Drei wichtige Entwicklungen, die vor allem in den letzten drei Jahrzehnten stattgefunden haben, machen dies möglich. Ersten werden in Unternehmen Daten zunehmend strukturierter und in größerer Menge erhoben, sodass Prozesse akkurat beschrieben werden können. Zweitens hat sich die Leistungsfähigkeit von Computern deutlich gesteigert, wodurch sie Rechenaufgaben sehr viel schneller als früher lösen können. Der dritte und wichtigste Punkt ist die Weiterentwicklung der mathematischen Methoden, die ein entscheidender Faktor für die Fortschritte im Bereich der KI sind.

Zwei Teilgebiete der Mathematik sind für viele derartige Problemstellung besonders wichtig: Statistik und mathematische Optimierung. Erstere ermöglicht es, eine Erwartung gegenüber der Zukunft auf Basis der Vergangenheit abzuleiten. *Wenn die Sonne die letzten hundert Jahre im Osten aufgegangen ist, wird sie mit großer Wahrscheinlichkeit auch morgen im Osten aufgehen.* Mathematische Optimierung errechnet demgegenüber in einer als festgelegt angenommenen Welt die beste Handlungsoption, um einen gewissen Nutzen zu maximieren. *Eine Photovoltaikanlage sollte so ausgerichtet sein, dass sie möglichst viel Sonne über den Tag hinweg einfangen kann. Dabei ist der Sonnenstand eine festgelegte, bekannte Größe.*

Entscheidungen in Supply Chains müssen allerdings häufig unter Unsicherheit getroffen werden: Auch, wenn die Entscheidenden die Rahmenbedingungen der Zukunft noch nicht genau kennen, müssen sie schon jetzt eine Entscheidung treffen, die den Nutzen in der Zukunft maximiert. Dieses Problem kann mithilfe der Kombination von Statistik und Optimierung gelöst werden. Statistik erlaubt es dabei, erwartete Rahmenbedingungen abzuschätzen, innerhalb derer dann mathematische Optimierung die beste Entscheidung findet. Gemäß Ihrer *Erwartung* (Statistik) des Sonnenverlaufs können Sie den *optimalen* (Optimierung) Neigungsgrad Ihrer Photovoltaikanlage ausrechnen.

Die Problemklassen, die Lieferketten in Unternehmen betreffen, sind der Frage, wie eine PV-Anlage ausgerichtet werden soll, häufig sehr ähnlich. Wang et al. (Wang et al. 2016) haben sieben Unternehmensbereiche identifiziert, in denen Anwendungen für datengetriebene Entscheidungsunterstützung erforscht wurden. Dabei lassen sich strategische, taktische und

operative Entscheidungen differenzieren. Diese Unternehmensbereiche umfassen auf strategischer Ebene weitsichtige Beschaffung, Design des Liefernetzwerks sowie Produktdesign und Entwicklung. Auf operativer bzw. taktischer Ebene nennen die Autoren Absatzplanung, Einkauf, Produktion, Lagerhaltung und Logistik.

Im Folgenden wird zwischen den Wissenschaftsfeldern „Data Science“ und „mathematischer Optimierung“ unterschieden. Beide Felder sind im Forschungsgebiet der **Künstlichen Intelligenz (KI)** einzuordnen. Die Trennung der beiden Felder anhand der Verwendung von datengetrieben bzw. modellbasierten Verfahren ist nicht immer eindeutig, jedoch fallen die meisten Lösungen für Data Science Problemstellungen in erstere Kategorie. Der Begriff Data Science ist in der Literatur nicht eindeutig definiert, wird aber oft als interdisziplinäres Forschungsfeld bezeichnet, welches statistische und technisch-informatische Verfahren vereint. Problemstellungen der Optimierung werden mit modellbasierten Verfahren gelöst.

1.1 ADA LOVELACE CENTER

Die fortschreitende Digitalisierung der Supply Chain und die Entwicklung immer neuer KI-Methoden fordert von Wissenschaftler*innen und Anwender*innen immer auf den neuesten Stand der Technik zu bleiben. Hierfür wurde das ADA Lovelace Center für Analytics, Daten und Anwendungen gegründet - ein KI-Kompetenznetzwerk bestehend aus dem Fraunhofer-Institut für Integrierte Schaltungen IIS, der Friedrich-Alexander-Universität Erlangen-Nürnberg (FAU) und der Ludwig-Maximilians-Universität München (LMU). Abbildung 1 zeigt den Aufbau und das Angebot des ADA Lovelace Centers als KI- Netzwerk, das Industrie und Forschung miteinander verbindet.

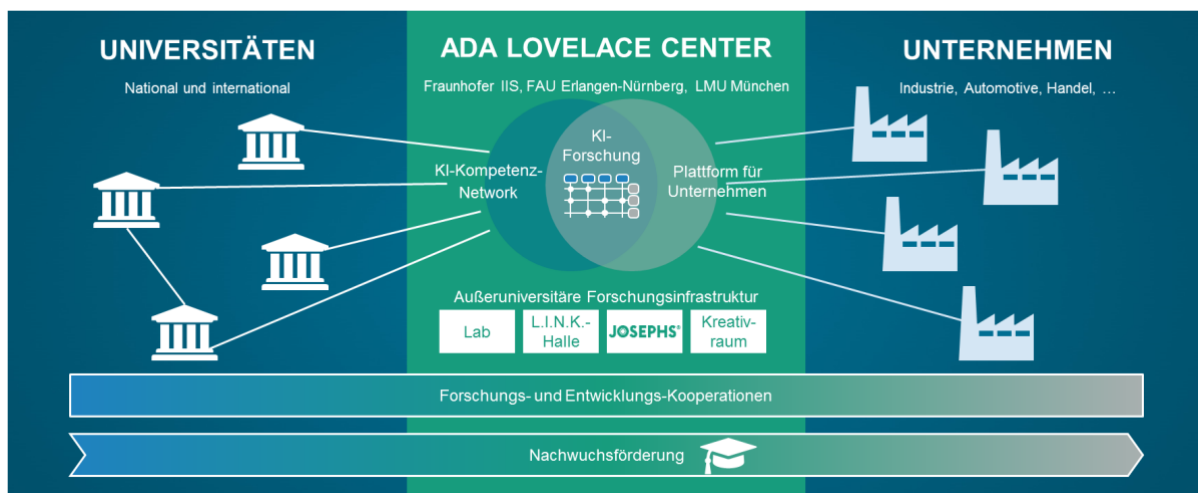


Abbildung 1: Aufbau und Angebot des ADA Lovelace Centers - Vernetzung von Industrie und Forschung

Gemeinsam erforschen die beteiligten Institute im ADA Lovelace Center die mathematischen Grundlagen der künstlichen Intelligenz, entwickeln performante KI-Verfahren für spezifische Anwendungen neu und bringen diese im Rahmen zahlreicher Industriekooperationen bis zur Praxisreife. Zudem stellt das ADA Lovelace Center eine Plattform für Unternehmen dar, um sich mit führenden KI-Forscher*innen und -Entwickler*innen zu vernetzen und vom Know-how von Wissenschaftler*innen nationaler und internationaler Universitäten zu profitieren. Die Forschungs- und Entwicklungsprojekte des ADA Lovelace Centers mit der Industrie reichen dabei von der Durchführung von Machbarkeitsstudien bis hin zur Anfertigung maßgeschneiderter Software-Prototypen für spezialisierte Aufgabenstellungen. Die Zusammenarbeit mit dem ADA Lovelace Center im Rahmen von Joint Labs ermöglichen den industriellen Partnern den Anschluss an state-of-the-art

KI-Expertise, leistungsfähige IT-Infrastruktur und KI-Software, Visualisierungstechnologien sowie hands-on Ausbildungsmodule. Außerdem unterstützt das ADA Lovelace Center die Migration von Studierenden und jungen Wissenschaftler*innen in die Wirtschaft. Dazu tragen die Betreuung von Abschlussarbeiten und Dissertationen im Rahmen von Industrieprojekten und das gemeinsame Arbeiten an Forschungsprogrammen der beteiligten Forschungseinrichtungen bei. Weiterhin erhalten KI-interessierte Unternehmen auf Veranstaltungen, über Wettbewerbe und Kontaktvermittlung direkten Zugang zu jungen Talenten. Ziel ist es, den Abstand zwischen der KI-Expertise der Nachwuchswissenschaftler*innen zu jener von erfahrenen IT-Expert*innen der Unternehmen und dem des mittleren Managements kontinuierlich zu verringern und so die Barrieren für die Nutzung von KI-Methoden in Unternehmen abzubauen.

Um stets auch auf dem globalen Stand des Wissens im Bereich der KI zu sein, hat das ADA Lovelace Center strategische Kooperationen mit international anerkannten Forschungseinrichtungen wie dem Center for Machine Learning des Georgia Institute of Technology in Atlanta, USA, und dem Riken Institute for Advanced Intelligence in Tokyo, Japan geschlossen. Die entwickelten Anwendungen im ADA Lovelace Center decken unter anderen die Schlüsselgebiete entlang der Supply Chain wie Produktion, Mobilität und Logistik ab.

1.2 KI-FRAGESTELLUNGEN IN DER SUPPLY CHAIN

Typische datengetriebene Fragestellungen entlang der Supply Chain können mit Methoden der Zeitreihenprognosen beantwortet werden. Oftmals soll ein Bedarf vorhergesagt werden, z. B. um darauf aufbauend Bestellentscheidungen zu treffen und so Lieferengpässe zu vermeiden. Auch für die Planung des Einsatzes von Personal und Ressourcen ist das Wissen über zukünftige Bedarfe ausschlaggebend. Des Weiteren können Preisprognosen eine große Unterstützung sein, um beispielsweise den Einkauf von Rohstoffen besser planen zu können. Im Bereich der Produktion ist die Analyse und Prognose von Durchlaufzeiten von Interesse. Außerdem stellt sich häufig die Frage, wann, wo und warum Fehler auftreten können. Hier können mit Data Science Methoden Abhängigkeiten z. B. zwischen verschiedenen Arbeitsschritten oder dem verwendeten Material und den aufgetretenen Fehlern aufgedeckt und analysiert werden. Ein weiterer Mehrwert von datengetriebenen Lösungen in der Supply Chain entsteht durch die Vorhersage von Maschinendefekten (Predictive Maintenance). So können Instandhaltungen rechtzeitig durchgeführt und Ausfallzeiten vermieden werden. Im Kapitel 2.1 werden einige ausgewählte Methoden aus dem Wissenschaftsfeld Data Science vorgestellt, welche Unternehmen dazu befähigen ihre Daten zu nutzen, um präzise Vorhersagen bezüglich Nachfragen, Preise und der Produktion zu machen. All dies führt zu einer radikalen Verbesserung der Effektivität und Effizienz der Logistik- und Produktionsprozesse.

Beispiele für mögliche Anwendungen von mathematischer Optimierung sind Planung von Personaleinsatz, Logistik- und Transportnetzwerken, Touren von Auslieferungsfahrzeugen und Produktionsplanung. Viele Optimierungsfragestellungen sind Planungsprobleme, bei denen entschieden werden muss, wie die Ressourcen eines Unternehmens bestmöglich eingesetzt werden. Der hohe händische Aufwand in der Planung, der in vielen Unternehmen oftmals noch vorherrscht, macht die Erstellung dieser Pläne kosten- und zeitintensiv. Zudem werden die Pläne typischerweise aus einer Menge lokaler Entscheidungen, also Teillösungen für einzelne Aspekte zu einer Gesamtlösung zusammengesetzt, was sich mit der Komplexität der Aufgabe erklären lässt. Die so entstandenen Lösungen verschwenden oft große Potenziale, die sich typischerweise aus Synergien zwischen den einzelnen zu lösenden Teilaufgaben ergeben. Mit Hilfe der mathematischen Optimierung lassen sich für solche Fragestellungen Softwarelösungen entwickeln. Somit können sich Planer*innen auf die größeren strategischen Fragestellungen konzentrieren. Durch die ganzheitliche Sicht der mathematischen Optimierung auf die Planungsaufgabe lassen sich in der Regel deutlich bessere Lösungen finden. Die in Kapitel 2.2 genannten Optimierungsmethoden versetzen

Planer*innen in die Lage, solche Pläne nach unterschiedlichen Zielkriterien effizient zu erstellen, vergleichen und bewerten. Dies ist besonders wichtig, wenn unvorhergesehene Ereignisse auftreten und schnell auf mögliche Änderungen der Eingangsdaten reagiert werden muss. Mit Hilfe der mathematischen Optimierung lässt sich so das wirtschaftliche Risiko für die ganze Produktion auf Grund der erhöhten Reaktionsgeschwindigkeit deutlich reduzieren. Auch Planspiele für verschiedene Möglichkeiten der strategischen Ausrichtung eines Unternehmens lassen sich so effizient durchführen.

Die angesprochene Verknüpfung von Data Science Methoden mit mathematischer Optimierung kommt dann zum Tragen, wenn die optimale Lösung von noch unbekannten Größen abhängt. Man spricht hier von präskriptiven Analysemethoden. Die optimale Bestückung eines Lagers, zum Beispiel, hängt natürlich von den zukünftigen Bestellungen der Kunden ab, die aber im Allgemeinen unbekannt sind. Können diese allerdings zuverlässig prognostiziert werden, kann darauf basierend die optimale Lösung gefunden werden. So ist es möglich, optimal vor auszuplanen.

Die Forschungsgruppen „Data Science“ und „Optimization“ der Fraunhofer-Arbeitsgruppe für Supply Chain Services des Fraunhofer-Instituts für Integrierte Schaltungen IIS erforschen datengetriebene und modelbasierte Lösungsverfahren entlang der Supply Chain. Dabei fokussieren sie sich zum einen auf Data Science Verfahren, die in Kapitel 2.1 vorgestellt werden, und zum anderen auf Verfahren der mathematischen Optimierung (Kapitel 2.2). Wie diese Verfahren in verschiedenen Industrie- und Forschungsprojekten konkret eingesetzt und kombiniert werden, ist in Kapitel 3 anhand einiger Beispiele beschrieben. Die Forschungsgruppen haben dabei den Anspruch, die Forschungsgrenze in der Anwendung stets weiter zu verschieben. In Kapitel 4 wird deswegen ein Ausblick auf die zukünftige Forschung gegeben.

2 ANWENDUNGSNAHE METHODEN DER GRUPPEN „DATA SCIENCE“ UND „OPTIMIZATION“ DER FRAUNHOFER-ARBEITSGRUPPE SCS DES FRAUNHOFER IIS

In diesem Kapitel werden die wichtigsten Methoden aus Data Science und mathematischer Optimierung, die in den Forschungsgruppen „Data Science“ und „Optimization“ verwendet werden, vorgestellt. Dabei soll ein möglichst breiter Überblick gegeben und die Methoden auch ohne mathematischen oder statistischen Hintergrund verständlich gemacht werden. Für einen tieferen Einblick sind die angegebenen Quellen empfohlen.

2.1 DATA SCIENCE

Data Science bezeichnet die Wissenschaft, Daten zu analysieren und anhand der gewonnen Erkenntnisse Fragestellungen zu beantworten. Im Zuge der Digitalisierung nimmt die Menge und Diversität der Daten zu. Zeitgleich entwickelt sich auch die Rechenleistung der Computer immer weiter und beschleunigt die Datenverarbeitung. Somit gewinnt Data Science immer mehr an Bedeutung. Um Daten auszuwerten, werden Methoden verschiedener verwandter Disziplinen verwendet und weiterentwickelt, wie z. B. der Statistik, des maschinellen Lernens, der Mathematik und der Informatik.

In der Anwendung werden Verfahren der **klassischen Statistik** und **maschinelles Lernen** (engl. machine learning, häufig mit ML abgekürzt) unterschieden. Bei statistischen Verfahren steht die Inferenz, also die Schlussfolgerungen, über das Modell im Fokus. Dabei werden i.A. die vorliegenden Daten als Stichprobe einer zugrundeliegenden Wahrscheinlichkeitsverteilung aufgefasst. In der Regel müssen hier also nur die Parameter der Verteilung geschätzt werden. Beim ML werden dagegen solche Annahmen im Allgemeinen nicht gemacht. Vielmehr ist hier das Ziel, nur anhand der Daten und ohne genauere Vorgaben Gesetzmäßigkeiten zu erkennen und daraus einen praktischen Nutzen zu ziehen. Inzwischen werden Ideen und Ansätze aus der klassischen Statistik und dem ML immer häufiger kombiniert, sodass diese Unterscheidung nicht immer sinnvoll ist.

Microsoft hat 2016 eine Darstellung der verschiedenen **Phasen eines Data-Science-Projekts** entwickelt, der sogenannte Team Data Science Process (TDSP) (What is the Team Data Science Process? 2016). In derselben Quelle findet sich auch eine grafische Darstellung dieses Prozesses. Dabei unterscheiden die Autor*innen fünf Phasen, die allerdings nicht chronologisch durchlaufen werden, sondern dynamisch verknüpft sind und iteriert werden können. In der ersten Phase *Business Understanding* werden im Austausch zwischen Data Scientist und Kunden das zu bearbeitende Problem definiert und die Ziele, Erwartungen und Risiken besprochen. Dies kann selten unabhängig von der zweiten Phase *Data Acquisition and Understanding* geschehen, in der die verfügbaren Daten untersucht und gegebenenfalls aufbereitet werden (**Preprocessing**). Dabei werden u. a. die folgenden Fragen beantwortet: Können verschiedene Datensätze sinnvoll zusammengefügt werden? Gibt es fehlende Werte und wie kann mit diesen umgegangen werden? Muss das Format der Variablen angepasst werden? Gibt es Fehler in den Daten und können diese herausgefiltert oder korrigiert werden?

In der Phase *Modeling* werden anhand der Erkenntnisse aus *Business Understanding* und *Data Acquisition and Understanding* sinnvolle Ansätze und Verfahren ausgewählt und anschließend Modelle geschätzt (mit statistischen Methoden) bzw. gelernt (mit ML).

Die Auswahl der eingesetzten Einflussvariablen, die sogenannte **Feature Selection** (Guyon und Elisseeff 2003), kann bereits im Rahmen des Preprocessings geschehen, also der Datenvorbereitung. Oftmals ist es aber sinnvoll, die Variablen danach auszuwählen, wie entscheidend ihr Einfluss im

jeweiligen Modell ist. Dies ist vor der Anwendung des Modells auf den Daten unbekannt. Feature Selection spielt dann eine besonders große Rolle, wenn viele Variablen, aber nur wenige Beobachtungen in den verfügbaren Daten enthalten sind. Bei der Verwendung aller Variablen kann es sonst schnell zu einer Überanpassung an die Daten kommen, sogenanntes **Overfitting**. Dies führt dazu, dass sich die Ergebnisse nicht auf die Grundgesamtheit außerhalb der analysierten Daten verallgemeinern lassen. Aber auch zur Interpretierbarkeit der Modelle trägt es bei, wenn die Variablen mit dem größten Einfluss identifiziert werden. Zudem können der rechnerische Aufwand und die Laufzeit verbessert werden, wenn nur eine kleinere Menge an tatsächlich relevanten Variablen verwendet wird anstelle eines umfassenden Modells mit allen verfügbaren Variablen.

Die gewählten Modelle, die sich u.a. in der grundlegenden Methodik, der Wahl der Parameter oder in der Auswahl und Verwendung der Einflussvariablen unterscheiden, gilt es im Anschluss zu vergleichen, um das beste Modell zu bestimmen. Aber was heißt „das beste Modell“?

Es gibt einige Möglichkeiten, die **Güte eines Modells** zu definieren und auch zu quantifizieren, damit Modelle verglichen werden können. Im Folgenden wird hierzu z. B. auf (James et al. 2017) verwiesen. Meistens wird die Anpassung des Modells an die zum Schätzen bzw. Lernen verwendeten Daten sowie das Verhalten des Modells auf bisher unbekannten Daten untersucht. Letzteres ist deshalb wichtig, da in vielen Fällen die Modelle für Aussagen über neue Daten genutzt werden sollen, wie beispielweise bei Klassifikationen. Es kann allerdings vorkommen, dass ein Modell die Daten der verwendeten Stichprobe perfekt beschreibt, darüber hinaus aber keinerlei Aussagekraft für unbekannte Daten besitzt. Dies muss bei der Wahl des besten Modells berücksichtigt werden. Häufig wird auch die Komplexität eines Modells miteinbezogen, da Modelle mit vielen Parameter*innen schwieriger zu interpretieren sind und mehr Rechenleistung erfordern, wobei letzteres mit der Entwicklung immer leistungsfähigerer Computer zusehends an Bedeutung verliert.

Die Anpassung statistischer oder ML Modelle an die Daten kann mithilfe der sogenannten **Likelihood** quantifiziert werden. Dafür wird angenommen, dass das untersuchte Modell die Wirklichkeit tatsächlich beschreibt. Unter dieser Annahme kann dann die Wahrscheinlichkeit berechnet werden, die vorliegenden Daten zu erhalten. Diese bedingte Wahrscheinlichkeit ist die Likelihood. Ist sie groß spricht das dafür, dass das betrachtete Modell zu den Daten passt, und umgekehrt. Um diese Anpassung an die Daten noch ins Verhältnis mit der Komplexität des Modells zu setzen, werden mithilfe der Likelihood sogenannte Informationskriterien, wie bspw. das Akaike Informationskriterium (AIC) oder das Bayes-Informationskriterium (BIC) berechnet. Ein anderer statistischer Ansatz ist das Bestimmtheitsmaß R^2 , das den Anteil der Varianz in den Daten, die das Modell erklärt, beschreibt.

Für Methoden ohne Verteilungsannahmen (wie häufig beim ML) kann die Likelihood allerdings nicht berechnet werden. Außerdem wird mit der Likelihood die Verallgemeinerbarkeit des Modells auf unbekannte Daten noch nicht getestet. Dieses Problem lösen **Kreuzvalidierungsverfahren** (engl. Cross-Validation). Hierbei werden Modelle nur auf Teilen der zur Verfügung stehenden Daten geschätzt bzw. trainiert und anschließend auf den restlichen getestet. Somit kann die Performance eines Modells auf unbekannten Daten simuliert werden. Je größer bzw. mehr die Fehler, die das Modell dabei auf den Testdaten macht, desto schlechter ist das Modell. Wie genau diese Fehler definiert sind und ob z. B. große Fehler schwerer gewichtet werden sollen als kleine Fehler, muss anhand des Anwendungsfalls entschieden werden. Eine Variante der Kreuzvalidierung ist, die Daten, auf denen trainiert bzw. getestet wird, mithilfe von **Bootstrapping**-Methoden zu bestimmen. Hierbei werden aus den vorliegenden Daten neue Stichproben mit Zurücklegen gezogen.

Ist nun das beste Modell gefunden, werden die Ergebnisse in der Phase *Deployment* an den Kunden übergeben. Akzeptiert der Kunde diese, ist auch die letzte Phase *Customer Acceptance* durchlaufen. Dies stellt allerdings nicht zwangsläufig das Ende des Projekts dar. Beispielsweise kann das Modell angepasst werden, wenn auch weiterhin aktuelle Daten mitberücksichtigt werden sollen.

Im Folgenden sollen nun die wichtigsten Methoden der Data Science kurz beschrieben werden. Erneut wird hier (James et al. 2017) für einen tiefergehenden Überblick empfohlen. Strukturiert sind

die Methoden danach, ob sie auf annotierten oder unannotierten Daten angewendet werden. Auch ihre Vor- und Nachteile sollen dargelegt werden. Hierbei ist allerdings zu beachten, dass dies nur sehr allgemeine Aussagen sind und im konkreten Anwendungsfall immer mehrere Methoden auf den Daten verglichen werden sollten, um die beste zu finden.

2.1.1 ÜBERWACHTES LERNEN

Eine Art von Daten, die vorliegen kann, sind beschriftete bzw. annotierte Daten (engl. labeled data). Wie der Name schon sagt ist bei annotierten Daten jedem Datenpunkt ein „Label“ zugeordnet. Von Interesse ist es, dieses Label für Daten vorherzusagen, bei denen es nicht bekannt ist. Klassische Fragestellungen, die sich aus annotierten Daten ergeben, sind also die **Klassifikation** (die Zuordnung zu diskreten Klassen) und die **Regression** (die Zuordnung zu Werten aus einem Wertebereich). Viele Methoden für annotierten Daten können mit geringen Anpassungen für beide Fragestellungen verwendet werden. Da man bei annotierten Daten überprüfen kann, ob ein Modell ein Label auf den zur Verfügung stehenden, bekannten Daten richtig zuordnet bzw. wie sehr sich das zugeordnete Label vom wahren unterscheidet, wird ML in diesem Gebiet auch **überwachtes Lernen** (engl. supervised learning) genannt.

Eine populäre Methode zur Klassifikation und Regression sind Entscheidungsbäumen (engl. **Decision Trees**). Dabei werden Regeln für Entscheidungen gelernt, die auch als Baum gut visualisiert und interpretiert werden können (Gordon et al. 1984), siehe Abbildung 2. Anhand der Features, d.h. der Eingabeparameter kann dieser Entscheidungsbaum im Anschluss durchlaufen werden und eine Kategorie bzw. ein Wert zugewiesen werden. Um die Gefahr von Overfitting zu reduzieren, kann der Baum „gestutzt“ werden (engl. Pruning). Allerdings sind Entscheidungsbäume relativ instabil, d.h. schon kleine Änderungen in den Daten können große Veränderungen im Aufbau des Baums bewirken. Um dem entgegen zu wirken, werden beim **Random Forest** mehrere unkorrelierte Entscheidungsbäume trainiert und die Prognose dann per Mehrheitsentscheid oder Mittelung zugewiesen (Breiman 2001). Random Forests sind effizient und sehr viel stabiler als einzelne Entscheidungsbäume. Allerdings geht hierdurch die einfache Interpretierbarkeit des Entscheidungsbaums verloren.

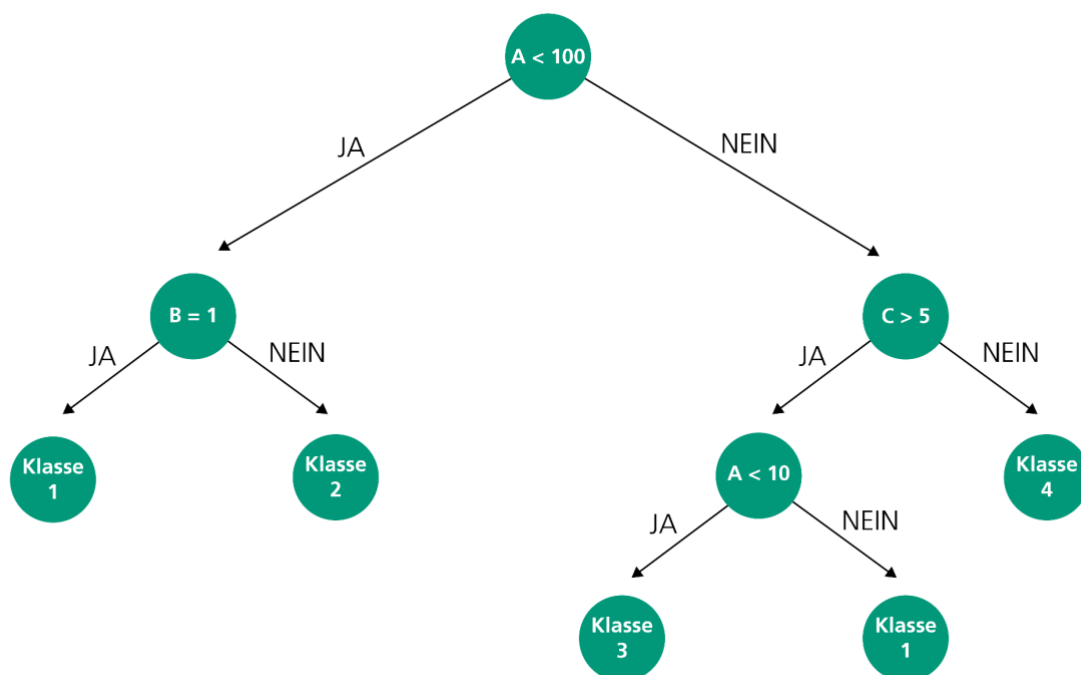


Abbildung 2: Graphische Darstellung eines Entscheidungsbaumes

Mit **Support Vector Machines (SVM)** wird eine weitere Methode bezeichnet. Die Idee dahinter ist, Datenpunkte so nach Klassen zu trennen, dass um die Klassengrenze herum ein möglichst breiter Streifen keine Datenpunkte enthält (siehe Abbildung 3 und (Cristianini und Shawe-Taylor 2012)). So können neue Datenpunkte einer Klasse zuverlässig zugewiesen werden, selbst wenn sie in diesem Streifen liegen. Um diese Grenzen zu bestimmen sind nicht alle Datenpunkte, sondern nur die sogenannten Stützvektoren (engl. support vectors) nötig. Da diese Trennung im Raum der Features im Allgemeinen nicht immer linear sinnvoll ist, wird mithilfe einer geeigneten Kernelfunktion das Problem in einen höherdimensionierten Raum und zurück transformiert. Die Wahl des Kernels und der Parameter sind allerdings entscheidend für die Genauigkeit der Methode. Bei großen Datenmengen ist außerdem der Rechenaufwand relativ groß. Dafür sind SVMs stabil und können Overfitting verhindern. Sind die Klassen separierbar, liefern SVMs häufig bessere Ergebnisse als andere Methoden. SVMs wurden ursprünglich für die Trennung zweier Klassen entwickelt, können allerdings auf multinomiale Klassifikationsprobleme und Regressionen erweitert werden.

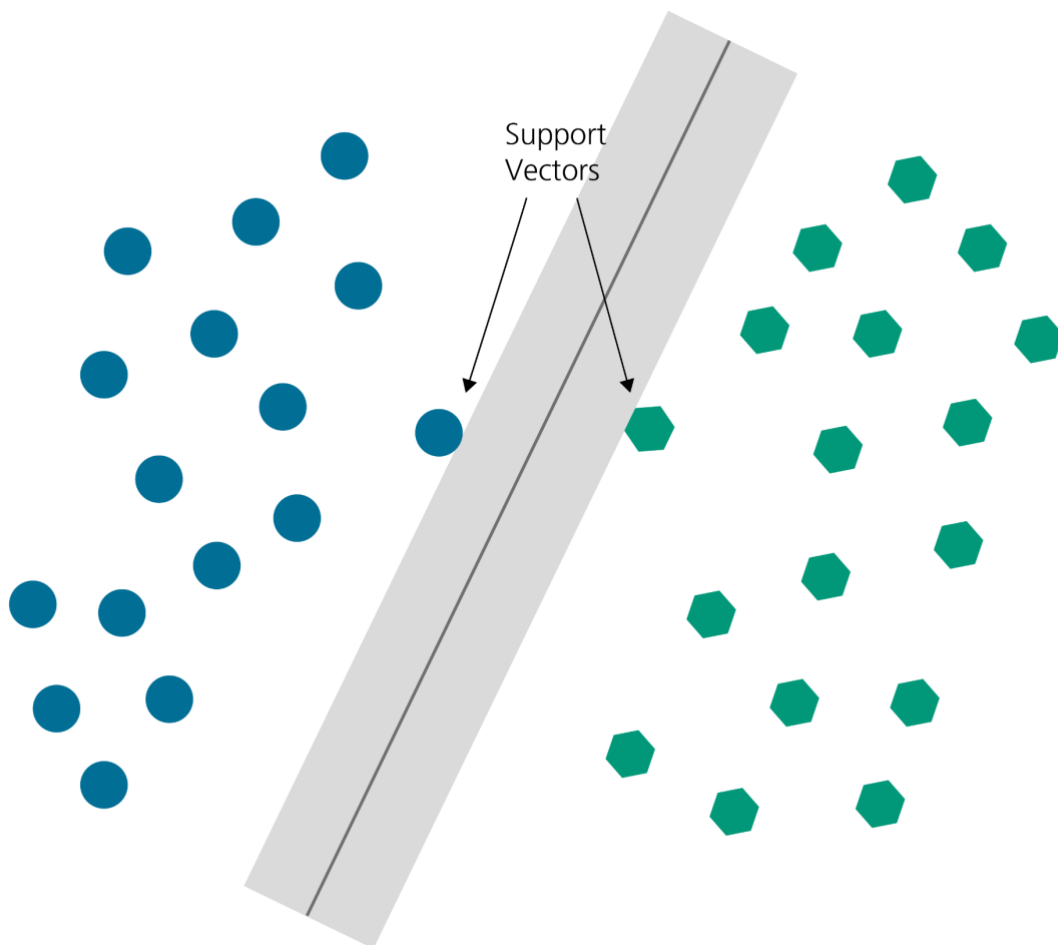


Abbildung 3: Graphische Darstellung Support Vector Machine

Eine weitere bekannte Methode sind **Neuronale Netze (NNs)**. Sie bestehen aus sogenannten künstlichen Neuronen, die in Schichten angeordnet und miteinander verbunden sind (Aggarwal 2018). Obligatorisch in einem klassischen Feed-Forward-Netz sind eine Input- und eine Output-Schicht, dazwischen können noch weitere, „versteckte“ Schichten (engl. hidden layers) verbunden sein (siehe Abbildung 4). Bei mehreren Zwischenschichten spricht man von tiefen neuronalen Netzen bzw. dem sogenannten Deep Learning. Jedes Neuron einer Schicht erhält die Outputs der Neuronen aus der letzten Schicht, die mit ihm verbunden sind. Aus diesen Outputs wird eine gewichtete Summe gebildet, die wiederum an eine Aktivierungsfunktion übergeben wird. Der Wert dieser Aktivierungsfunktion stellt dann den Output des Neurons dar und wird an die nächste Schicht weitergegeben. Diese Gewichte, mit denen die Neuronen-Outputs in die Outputs der Neuronen der

nächsten Schicht einfließen, sind die Parameter, die während der Trainingsphase gelernt werden. Die beschriebene Architektur kann allerdings noch auf verschiedene Weisen an Problemstellungen angepasst werden, um bessere Ergebnisse zu erzielen. Am Ende dieses Kapitels werden Architekturen vorgestellt, die sich besonders für die Analyse von Zeitreihen eignen.

In Neuronale Netze benötigen viele Trainingsdatenpunkte und eine große Rechenleistung. Steht dies allerdings zur Verfügung, funktionieren NNs häufig besser als andere Methoden und benötigen zudem keine Verteilungsannahmen. Deshalb werden sie z. B. in der Bilderkennung eingesetzt. Sollen am Ende Schlüsse über die Art der Abhängigkeit zwischen Input- und Outputvariablen aus den geschätzten Parametern des Modelles gezogen werden, eignen sich Neuronale Netze hingegen nicht.

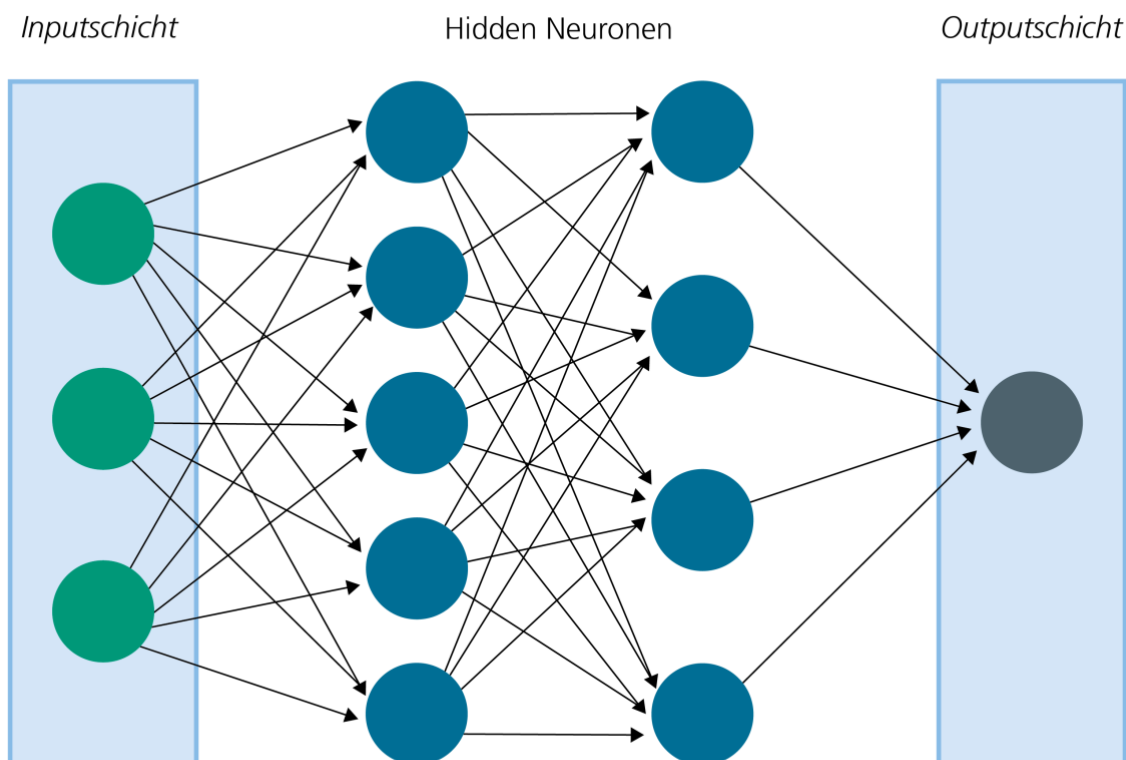


Abbildung 4: Graphische Darstellung eines neuronalen Netzes mit zwei versteckten Schichten

Während man die bisher vorgestellten Methoden dem ML zuordnen würde, entstammen **generalisierte lineare Modelle** (engl. generalized linear models, GLMs) der Statistik. Wie der Name schon vermuten lässt, stellen sie eine Verallgemeinerung der linearen Regression dar (Nelder und Wedderburn 1972). Bei dieser wird angenommen, dass der Erwartungswert der zu erklärenden („abhängigen“) Variable als Linearkombination der Features beschrieben werden kann und die abhängige Variable normalverteilt ist. Häufig sind allerdings Zielvariablen von Interesse, bei denen diese Annahme nicht sinnvoll ist, wie z. B. binäre Ja-Nein-Entscheidungen oder Preise, die keine negativen Werte annehmen können. Um dieses Problem zu lösen, wird die abhängige Variable mithilfe einer sogenannten Link-Funktion mit den linearen Prädiktoren verknüpft. Außerdem können auch andere Verteilungen als die Normalverteilung für die abhängige Variable angenommen werden. Dies verleiht GLMs eine große Flexibilität. Zudem können durch die Verteilungsannahmen leicht Aussagen zu Prognosegenauigkeit, Konfidenzintervallen der Parameter sowie Art und Größe des Einflusses der Prädiktoren getroffen werden. Sind diese Annahmen allerdings irrtümlich getroffen, sind auch die Ergebnisse der Regressionsanalyse nicht verlässlich. Außerdem liefern ML- Methoden

in der Praxis häufig genauere Ergebnisse als GLMs, zumindest wenn eine relativ große Menge an Daten zum Lernen zur Verfügung steht. GLMs bieten sich also dann an, wenn die Verteilungsannahmen gerechtfertigt erscheinen und die Interpretation der Zusammenhänge im Fokus steht.

Als letztes sollen an dieser Stelle **Bayes'sche Netze** vorgestellt werden (Pearl 1985). Diese sind eine Methode, um die gemeinsame Verteilung verschiedener Zufallsvariablen kompakt mithilfe eines gerichteten Graphen darzustellen (siehe Abbildung 5). Dabei entsprechen die Knoten den Zufallsvariablen und die Kanten bedingten Wahrscheinlichkeiten. Die Struktur des Graphen kann sowohl angenommen werden, indem Expertenwissen herangezogen wird, als auch aus den Daten gelernt werden. Wird sie angenommen, werden zusätzlich die Parameter der bedingten Verteilungen geschätzt. Bayes'sche Netze werden häufig als Darstellung kausaler Zusammenhänge interpretiert und zur Ursachenanalyse verwendet. Dabei wird der Satz von Bayes genutzt, um zu Beobachtungen die Wahrscheinlichkeiten der jeweiligen möglichen Ursachen zu bestimmen. Im Gegensatz zu GLMs kann auch eine Wahrscheinlichkeitsverteilung über die zu schätzenden Parameter (nicht nur über ihre Schätzer) ausgegeben und somit die Genauigkeit intuitiv interpretiert werden. Ein weiterer Vorteil von Bayes'schen Netzen ist, dass sowohl fehlende Daten als auch kleine Datensätze i.A. keine Probleme darstellen. Der größte Nachteil von Bayes'schen Netzen ist allerdings der hohe Rechenaufwand, insbesondere wenn die Struktur des Netzes aus den Daten gelernt werden soll. Außerdem kann es schwierig sein, Expertenwissen in die Modelldarstellung zu überführen und sinnvolle Annahmen zu treffen, wie z. B. die Wahl der A-Priori-Verteilungen.

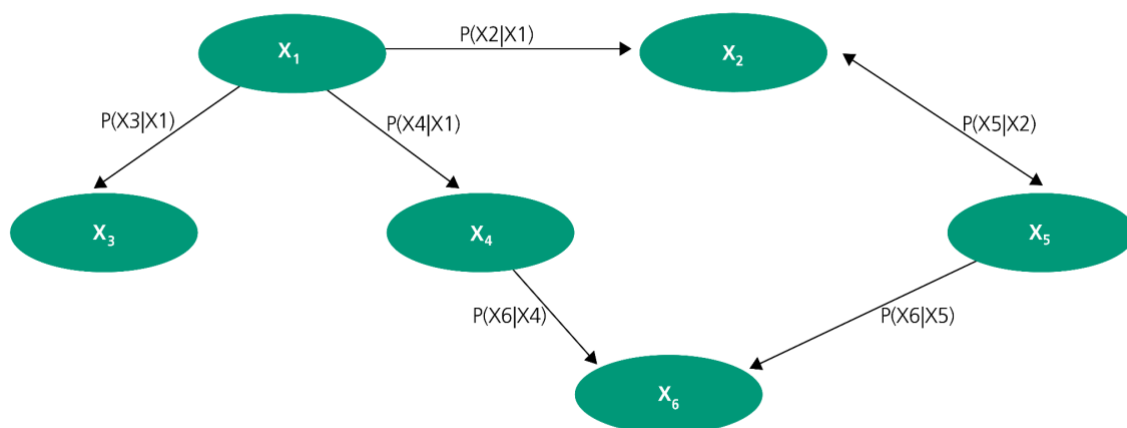


Abbildung 5: Graphische Darstellung eines Bayes'schen Netzes

Viele der bisher beschriebenen Methoden können auch zur **Prognose von Zeitreihen** benutzt werden. Von Zeitreihen spricht man, wenn eine Größe (wie beispielsweise der Absatz eines Produkts) über die Zeit beobachtet wird. In der Logistik beschreiben Zeitreihendaten häufig **dynamische Systeme**, also Systeme, bei denen sich der folgende Zustand aus dem vorherigen ableitet. Deshalb gibt es für Zeitreihendaten spezifische Methoden, die die Abhängigkeiten der einzelnen Beobachtungen voneinander berücksichtigen und die Struktur dynamischer Systeme inhärent in der Modellstruktur abbilden. Zudem gibt es bei Zeitreihen wiederkehrende Muster, wie beispielsweise Saisonalität, welche explizit modelliert werden können. In klassischen Methoden wird dabei häufig lediglich eine, sich verändernde Größe betrachtet – eine starke Vereinfachung, da dynamische Systeme in der Realität meistens aus verschiedenen wechselwirkenden Faktoren bestehen. Diese

Wechselwirkung kann lokal in Näherung linear beschrieben werden. Durch die Berücksichtigung auch nichtlinearer Effekte, erhält man meistens eine akkuratere Modellierung. So können beispielsweise Sättigungseffekte in Märkten oder zur Neige gehende Rohstoffe zu nichtlinearen Zusammenhängen führen.

Während einige Methoden zur Zeitreihenprognose bereits in verschiedenen ERP-Systemen integriert sind (beispielsweise SAP-APO), erfordern andere Methoden eine genauere Anpassung an die Problemstellung.

Die einfachste Idee zur Zeitreihenprognose ist, den letzten bekannten Wert als Prognose für die Zukunft zu verwenden. Dieses Verfahren wird „**naive Prognose**“ genannt und kann grundsätzlich als Vergleichsmodell für komplexere Modelle dienen. Im nächsten Schritt können zusätzlich weitere Werte aus der Vergangenheit berücksichtigt werden, von denen dann der gewichtete Mittelwert gebildet wird.

Intuitiv ergibt es dabei Sinn, den Werten, welche weniger weit zurückliegen, eine höhere Bedeutung beizumessen. Die Verfahren der **exponentiellen Glättung** implementieren diese Idee. Die Gewichtung der letzten Zeitschritte nimmt dabei exponentiell ab, je weiter zurück der beobachtete Wert liegt. Das Verfahren wurde ursprünglich von (Brown 1959) publiziert. Das darauf aufbauende **Holt-Verfahren** (Holt 2004)¹ kann dabei Trends berücksichtigen und das **Holt-Winters-Verfahren** (Winters 1960) zusätzlich Saisonalitäten.

Auch das **ARMA-Modell** (Box und Jenkins 1970) beruht auf der Idee, dass sich der zukünftige Datenpunkte aus einer gewichteten Summe vergangener Datenpunkte berechnen lässt. Darüber hinaus wird bei dieser Modellklasse angenommen, dass es Variablen gibt, die nicht beobachtet wurden. Deren Einfluss wird mithilfe des vergangenen Prognosefehlers geschätzt. Erweiterungen für Zeitreihen mit Trend bzw. Saisonalität heißen **ARIMA** bzw. **SARIMA**. Darüber hinaus wurde für multivariate Zeitreihenprognose die Erweiterung **ARIMAX** entwickelt.

Die bisher beschriebenen Verfahren erlauben beim Bau von Prognosemodellen nur geringe Freiheitsgrade. Demgegenüber können in Neuronalen Netzen problemspezifische Muster abgebildet werden. Speziell für die Prognose von Zeitreihen geeignete NNs sind beispielsweise **Rekurrente Neuronale Netze**. Sie basieren auf der Idee, dass der Zustand von Morgen vom heutigen Zustand abhängt. Im Rahmen dieser Annahme adressieren **Long Short-Term Memory-Netzwerke** die Frage, welche weit in der Vergangenheit liegenden Werte trotzdem Relevanz für die Zukunft haben (Hochreiter und Schmidhuber 1997). **Error Correction Neural Networks** nutzen ähnlich wie ARMA Modelle den Fehler der Vorperiode, um die nächste Periode zu prognostizieren (Zimmermann et al. 2002). Die Struktur eines ECNNs wird in Abbildung 6 abgebildet. Weitere, komplexere Architekturen umfassen **Historical Consistent Neural Network-Modelle (HCNN)**, welche keine klassische Input-Output-Relation annehmen, sondern stattdessen alle Wechselwirkungen des dynamischen Systems implizit lernen (Zimmermann et al. 2011). **Causal-Retrocausal Neural Networks (CRNN)** berücksichtigen Kausalitäten, welche von der Vergangenheit in die Zukunft gerichtet sind, genauso wie Kausalitäten, die Verbindungen von der Zukunft in die Vergangenheit herstellen (Zimmermann et al. 2012). Sie sind deshalb besonders gut geeignet, um Märkte zu beschreiben.

¹ Bei der Quelle handelt es sich um einen Nachdruck. Die ursprüngliche Version ist bereits 1957 erschienen.

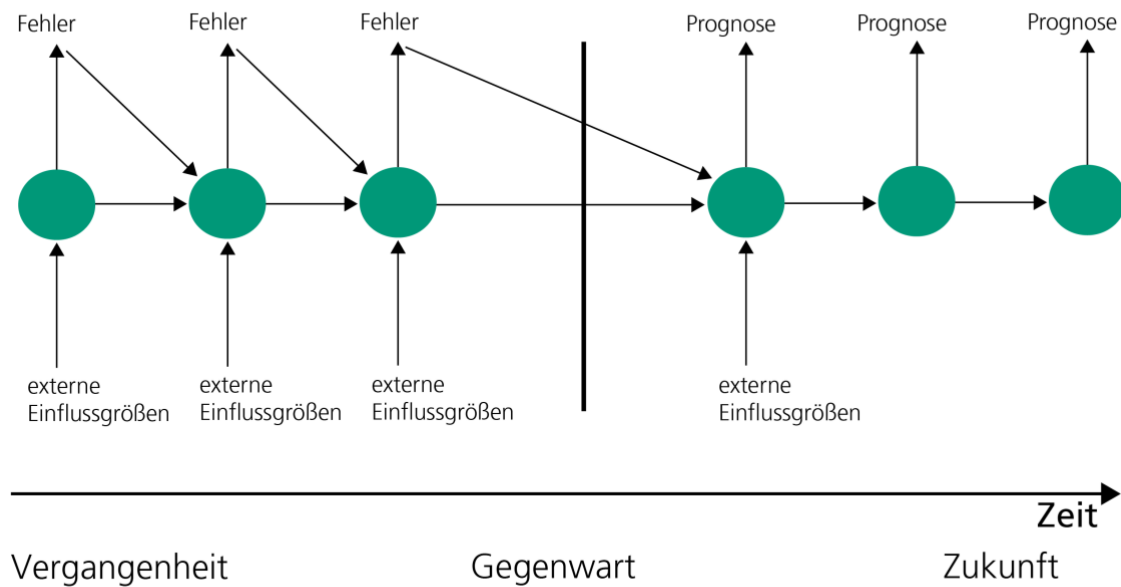


Abbildung 6: Graphische Darstellung eines Error Correction Neural Networks, bei dem der Fehler aus dem vorangegangenen Zeitschritt für die Vorhersage des nächsten Zeitschrittes verwendet wird.

2.1.2 UNÜBERWACHTES LERNEN

Eine zweite Kategorie von Daten sind die nicht annotierte Daten (engl. unlabeled data). Im Gegensatz zu annotierte Daten haben die Datenpunkte keine bekannte Zielgröße bzw. sind keiner Kategorie zugeordnet. Somit können Data Analytics Methoden nicht direkt evaluiert werden, da es aufgrund des fehlenden Labels keinen Vergleich zwischen Ausgabe des Modells und Zielgröße gibt. In diesem Zusammenhang spricht man von **unüberwachtem Lernen** (engl. unsupervised learning). Interessante Fragestellungen sind hier beispielsweise mögliche Gruppierungen von Datenpunkten oder das Auffinden von Mustern in Datensätzen. Einige Methoden können auch verwendet werden, um die ursprüngliche Dimension der Daten zu reduzieren, um weitere Analysen zu vereinfachen.

Ein häufig verwendeter Ansatz ist die **Clusteranalyse**. Sie bietet die Möglichkeit, ähnliche Datenpunkte zusammenzufassen (siehe Abbildung 7). Dadurch können Zusammenhänge entdeckt und Beobachtungen in Kategorien eingeteilt werden. Im Supply Chain Kontext wird dies beispielsweise auf Stammdaten angewandt werden, um Ähnlichkeiten in Artikeln zu identifizieren.

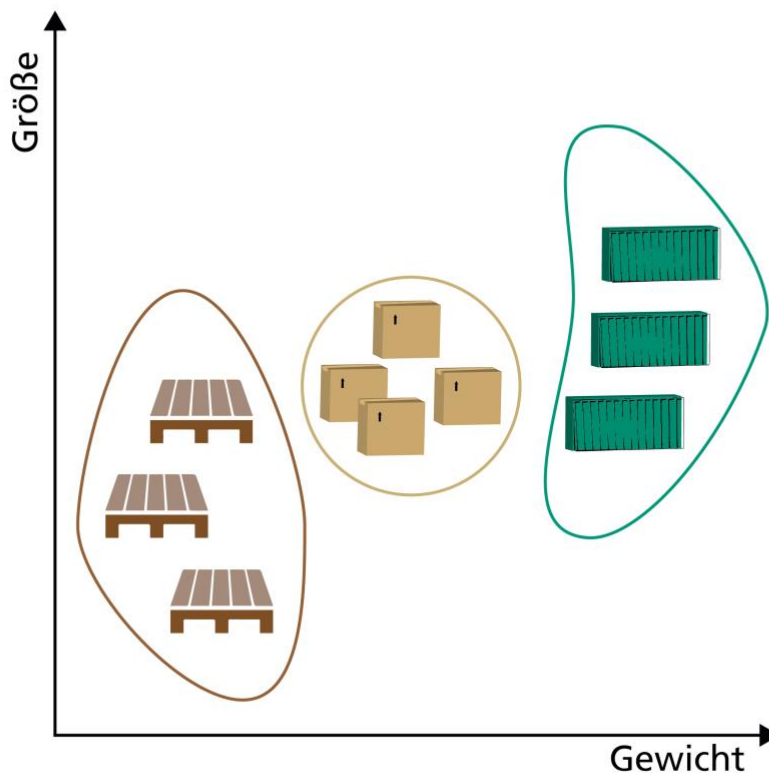


Abbildung 7: Graphische Darstellung einer Clusteranalyse

Bei der Clusteranalyse können verschiedene Herangehensweisen unterschieden werden. Zum einen existieren Algorithmen, bei denen die Anzahl der Cluster im Datensatz a priori vom Anwender festgelegt werden muss. Ein Beispiel ist der **k-means Clusteralgorithmus**, bei dem k Clusterzentren so im Datensatzraum positioniert werden, dass die Summe der Abstände von jedem Datenpunkt zu seinem nächstgelegenen Clusterzentrum minimal wird (Lloyd 1982). Dieser Algorithmus bietet viele Vorteile und ist unter anderem schnell und robust. Allerdings hängt die Performance auch stark von der gewählten Anzahl der Clusterzentren ab. Werden zu wenige gewählt, können verschiedene Gruppierungen nicht richtig getrennt werden. Sind es zu viele, werden einzelne Gruppen fälschlicherweise aufgeteilt. Das Problem, die Anzahl der Cluster (k) im Voraus per Hand wählen zu müssen, kann mit der zweiten Herangehensweise, der **hierarchischen Clusteranalyse**, umgangen werden (Johnson 1967). Unterschieden wird hier zwischen divisiven und agglomerativen Verfahren. Bei divisiven Algorithmen wird zu Beginn die gesamte Datenmenge als ein Cluster betrachtet und dieses schrittweise immer in kleinere Partitionen aufgeteilt. Im Gegensatz dazu wird in agglomerativen Methoden der umgekehrte Weg gegangen und jeder Datenpunkt als eigenes Cluster verstanden. Danach werden iterativ einzelne Cluster zu einer größeren Gruppe zusammengefügt. Bei diesen Verfahren kann die gewünschte Anzahl von Clustern im Nachhinein gewählt werden. Dies kann allerdings auch als Nachteil verstanden werden, da die Algorithmen eine sehr große Anzahl von möglichen Partitionierungen liefern und final daraus entschieden werden muss, welche Aufteilung die passendste ist.

Weder vor noch nach Ausführung des Algorithmus muss die Clusteranzahl bei der **DBSCAN** (Density-Based Spatial Clustering of Applications with Noise) Clusteranalyse gewählt werden (Ester et al. 1996). Sie bietet die Möglichkeit, die Bestimmung der Clusteranzahl dem Algorithmus selbst zu überlassen. Es wird die Annahme getroffen, dass in einem Cluster und dessen Rand die Dichte der

Datenpunkte höher ist als außerhalb. Somit wird je nach Dichte (Anzahl von Datenpunkten in einer gewissen Umgebung) entschieden, ob die Datenpunkte im selben Cluster liegen sollen. Ein weiterer Vorteil der Methode ist, dass der Algorithmus in der Lage ist, einzelne Ausreißer auch als solche zu identifizieren und keinem Cluster zuzuordnen. Dabei besteht das Zentrum eines Clusters immer aus denselben Punkten. Willkürlichkeit durch zufällige Startbedingungen ist hiermit ausgeschlossen. Nur Elemente, die zu zwei Clustern gehören könnten, werden zufällig zu einem der beiden zugeordnet.

Bei Zeitreihen hingegen können mit Hilfe des **Functional Fisher-Expectation-Maximization** (funFEM) Algorithmus unterschiedliche Zeitreihen mit einem ähnlichen zeitlichen Verlauf gruppiert werden (Bouveyron et al. 2015). Dafür werden zuerst diskrete Zeitdaten, beispielsweise mit Hilfe von Interpolation, in stetige Kurven transformiert. Danach werden durch eine Erweiterung des EM-Algorithmus die entsprechenden Cluster bestimmt. Zusätzlich kann durch eine Überlagerung von mehreren Basis-Funktionen (z. B. Gauß-Kurven) eine mittlere Zeitreihe aus den Elementen in einem Cluster bestimmt werden. Diese Mittelwertkurven können als Cluster-Zentren verstanden werden und geben einen Überblick über die verschiedenen Muster von Zeitreihen.

Ein anderer Anwendungsbereich von Methoden des unüberwachten Lernens ist die Dimensionsreduktion von Datensätzen. Dies kann auch als Aufbereiten der Daten verstanden werden und als Werkzeug zur **Feature Selection** (siehe oben) genutzt werden. Oft bestehen Datensätze aus vielen hochkorrelierten Features, was zu unnötig hohen Dimensionen in der Analyse führt. Teilweise besteht die Möglichkeit, die Anzahl der Dimensionen zu verringern und die Daten mit weniger Eigenschaften zu beschreiben, ohne dadurch relevante Informationen zu verlieren. Die **Hauptkomponentenanalyse** (engl. Principal Component Analysis) transformiert die vorhandenen Features mit Hilfe von Linearkombinationen in Hauptkomponenten (engl.: principal components) (Wold et al. 1987). Diese sind untereinander unkorreliert und sind nach Relevanz angeordnet. Somit beschreibt die erste Hauptkomponente die größte Varianz im Datensatz, die zweite die zweitgrößte Varianz, die von der ersten Komponente nicht erklärt wird, usw. Somit kann ein Großteil des Datensatzes mit den ersten Hauptkomponenten beschrieben werden und die letzten Hauptkomponenten gegebenenfalls vernachlässigt werden, da mit ihnen kein oder nur sehr wenig Informationsgewinn einhergeht. Durch das Vernachlässigen der irrelevanten Hauptkomponenten entsteht eine Reduktion der Dimension. Dies kann zu besseren Ergebnissen bei weiteren Analyseschritten führen (siehe Abbildung 8).

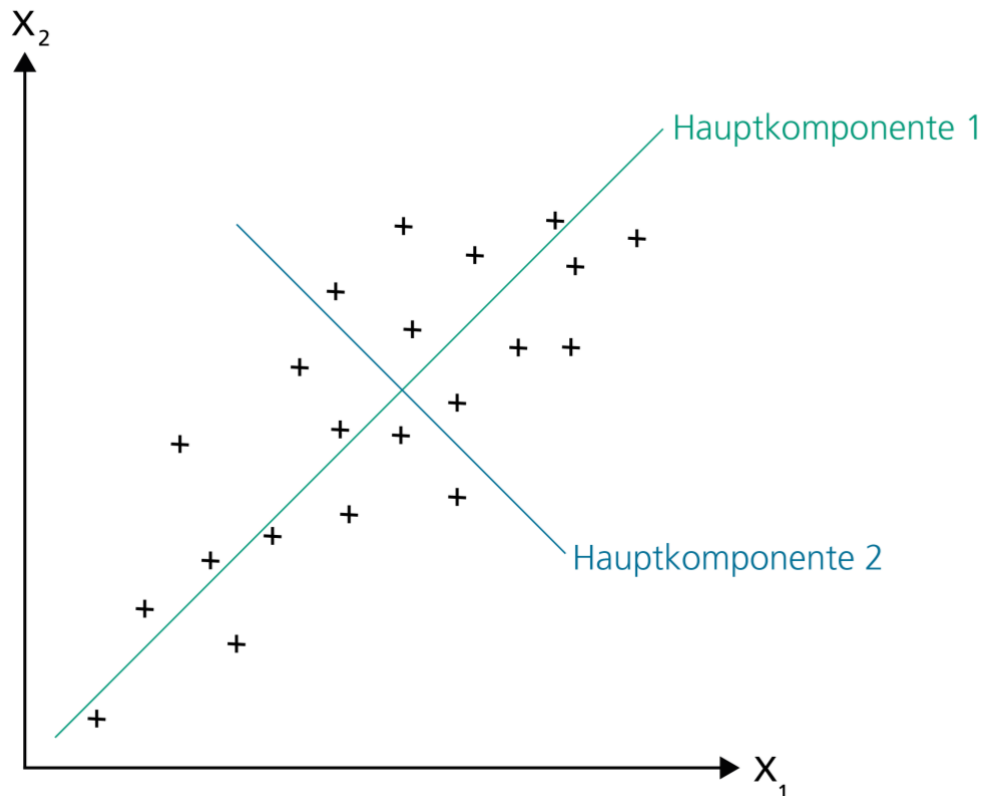


Abbildung 8: Graphische Darstellung einer Hauptachsentransformation. Die Hauptkomponenten sind orthogonal zueinander und damit unkorreliert.

2.2 MATHEMATISCHE OPTIMIERUNG

Während der umgangssprachliche Gebrauch des Wortes „Optimierung“ oft benutzt wird für „etwas besser machen“ oder, in frühen Phasen des Problemverständnisses, schlicht für „etwas anders machen“, meint die mathematische Optimierung das **gezielte Suchen einer optimalen Lösung für eine wohldefinierte Fragestellung**. Dazu werden die im Unternehmen gültigen Planungsregeln und Zielvorgaben in ein mathematisches Modell übersetzt, das alle denkbaren Lösungen für die gestellte Aufgabe berücksichtigt. Zur Lösung dieses Modells ist dann einerseits die Verfügbarkeit von Daten notwendig, die als Entscheidungsgrundlage betrachtet werden (z. B. Daten über die Rohstoffverfügbarkeit oder die erwartete Kundennachfrage). Andererseits werden leistungsfähige Algorithmen benötigt, welche das Optimierungsproblem innerhalb der geforderten Antwortzeit lösen können.

Ein mathematisches Optimierungsmodell wie oben beschrieben besteht aus einer oder mehreren Zielkriterien (sog. **Zielfunktionen**), **Entscheidungsvariablen** und **Nebenbedingungen** und richtet sich meist nach einer konkreten Planungsentscheidung, die im Unternehmen getroffen werden muss. Die Entscheidungsvariablen beschreiben die Stellschrauben der Planung. Sie repräsentieren oft kleine, einzelne Entscheidungen, die zusammengesetzt den ganzen Plan für die Fragestellung ergeben. Die Nebenbedingungen beschreiben betriebliche, gesetzliche oder weitere Regeln, die für einen gültigen Gesamtplan gelten müssen. Insbesondere schließen sie unerwünschte Lösungen und Effekte aus und beschreiben das Zusammenspiel und die gegenseitige Abhängigkeit der unterschiedlichen atomaren Entscheidungen. Die Zielfunktion schließlich bewertet eine Lösung nach fest definierten Kriterien und macht so verschiedene Lösungen vergleichbar: Lösung A ist besser, schlechter oder gleich gut wie Lösung B. Ein mathematisches Optimierungsverfahren ist auf Basis eines solchen Modells in der Lage, nach einer bestmöglichen Lösung für das gestellte Problem zu suchen.

Während sich die grundsätzliche Richtung des Optimierungsproblems an der Planungsfragestellung orientiert, haben die verfügbaren Daten oft einen entscheidenden Einfluss auf die Detailgestaltung der konkreten Modellierung. Die Verfügbarkeit, Qualität und Vollständigkeit beeinflussen einerseits die Geschwindigkeit, mit der eine Softwarelösungen entwickelt werden kann, andererseits entscheiden sie maßgeblich, wie vertrauenswürdig die Lösung des Optimierungsproblems ist. In der Praxis ist das Fehlen von entscheidungsrelevanten und gut aufbereiteten Daten oft eines der größten Hindernisse. Eine saubere Datenhaltung ist daher maßgeblich für den Erfolg eines Optimierungsprojektes mitverantwortlich.

Sind Daten in ausreichender Menge und Qualität für das Optimierungsproblem vorhanden, wird ein **Optimierungsalgorithmus** benötigt, um das Modell zu lösen. Nicht alle Optimierungsmodelle sind auf Anhieb in der Praxis lösbar. Daher ist es oftmals sinnvoll Experten für mathematische Optimierung hinzuzuziehen, die bei der Modellierung als auch bei der Auswahl oder, falls notwendig, dem maßgeschneiderten Entwurf von Lösungsverfahren beraten können.

Die Qualität und die Effizienz des gewählten Lösungsverfahrens sind dabei wichtige Aspekte. Das Feld der mathematischen Optimierung hat sich in den letzten 100 Jahren rasant weiterentwickelt und Jahr für Jahr werden die Verfahren leistungsfähiger. Es können immer kompliziertere Modelle gelöst werden, welche vor Jahren noch für unlösbar gehalten wurden.

Im Folgenden werden an Hand des weithin bekannten **Travelling-Salesman-Problems** (TSP, zu Deutsch: dem Problem des Handlungsreisenden) verschiedene Lösungsalgorithmen vorgestellt und deren Vor- und Nachteile diskutiert. Das Problem entsteht aus der typischen Fragestellung einer Tourenplanung: Es geht dabei darum, eine Rundreise zu verschiedenen Städten zu planen, sodass der Handlungsreisende bzw. der LKW am Ende wieder am Ausgangspunkt (dem Heimatort bzw. dem Depot) ankommt und eine minimale Strecke zurückgelegt hat. Während 1954 nur eine Instanz mit 49 Städten optimal gelöst werden konnte, wurde 2006 eine Instanz mit 1.9 Millionen Ort auf der Welt beweisbar fast optimal gelöst (bis auf 0,01% genau). Dies zeigt die enorme Weiterentwicklung der mathematischen Methoden. Eine Einführung zum TSP findet sich in (Applegate et al. 2006; Padberg und Rinaldi 1987).

2.2.1 HEURISTIKEN

Die simpelste Methode, um das TSP zu lösen, ist die **Brute-Force-Methode**. Bei dieser wird die Gesamtheit aller möglichen Touren aufgezählt, für jede einzelne Tour die Länge berechnet und am Ende die Beste ausgewählt. Dieses Verfahren stößt bereits bei verhältnismäßig kleinen Größenordnungen an seine Grenzen. Der Grund dafür ist das, was in der Mathematik die kombinatorische Explosion genannt wird. Mit steigender Anzahl an Städten, die besucht werden sollen, wächst die Anzahl der Touren exponentiell an. Während bei 5 Städten noch eine überschaubare Anzahl von 12 Touren möglich sind (äquivalente Lösungen zusammengefasst), sind es bei 30 Städten schon mehr Touren als es Sterne im Universum gibt. Selbst wenn man 1 Millionen Touren pro Sekunde berechnen könnte, würde man zu Lebzeiten keine optimale Lösung bekommen. Auch ein denkbarer schnellerer Computer bringt hier keinen nennenswerten Fortschritt. In der Praxis müsste also nach einer festen Zeit der Brute-Force-Ansatz abgebrochen werden, ohne eine Aussage über die Qualität der erhaltenen Lösung im Vergleich zur optimalen Lösung zu erhalten.

Eine zweite Klasse von Methoden sind sogenannte **Greedy-Heuristiken** (vom Engl. greedy: „gierig“, „gefräßig“). Sie bringen die Entscheidungen in eine Reihenfolge und wählen dann immer die lokal beste Entscheidung. Diese lokalen Entscheidungen sind oft durch die Zielfunktion oder „best practice“ Regeln motiviert. Eine der bekanntesten Vertreter dieser Klasse ist die **Nearest-Neighbour-Heuristik**. Bei ihr konstruiert man ausgehend vom Startort die Lösung schrittweise, indem man immer die zum aktuellen Ort nächstgelegene Stadt als nächstes besucht und am Schluss wieder zum Ausgangsort zurückkehrt. Während man durch diese Verfahren schnell eine Lösung erhält, sind in der Praxis die Lösungen oft deutlich suboptimal. Die Gründe dafür sind einerseits die zu lokale Sicht des

Algorithmus auf die getroffenen Entscheidungen und andererseits das fehlende Gesamtbild für die Interaktionen zwischen diesen Entscheidungen.

Eine dritte Klasse von Methoden sind **lokale Verbesserungsheuristiken** (LVH). Im Gegensatz zur Greedy-Heuristik, die eine zulässige Lösung für das Optimierungsproblem aufbaut, starten LVH bereits mit einer zulässigen Lösung und versuchen diese dann durch lokale Veränderungen zu verbessern. Ein bekannter Vertreter für das TSP ist der **2-opt Algorithmus**. Dieser iteriert über jedes Paar von Streckenabschnitten zwischen zwei Städten und überprüft ob die 4 betrachteten Städte sich kürzer verbinden lassen und tauscht, falls dies der Fall ist, die Strecken aus. Dieses Verfahren wird auch gerne kombiniert mit einer Greedy-Heuristik, um die daraus gewonnene Startlösung zu verbessern. Weitere bekannte Verfahren dieser Klasse sind Simulated Annealing (Peter J. M. van Laarhoven und Emile H. L. Aarts 1987), Tabu Search, Genetic Algorithms (Davis 1991) und Ant Colony Optimization (Dorigo et al. 1999).

Eine Übersicht über die bekanntesten Heuristiken findet sich in (Rego et al. 2011). Die vorgestellten heuristischen Verfahren haben den Vorteil, dass sie einfach zu implementieren und benutzen sind. Sie sind auch oft schon in Programmbibliotheken implementiert. Viele generische Heuristiken haben Namen aus der Natur, da sie von natürlichen Phänomenen inspiriert sind, wie z. B. der Ameisenalgorithmus oder genetische Algorithmen. Deren Funktionsweise ist in ihrer Anschaulichkeit für das Management typischerweise leicht verständlich und ihr Einsatz daher leicht zu rechtfertigen. Die Performance dieser generischen Algorithmen hinkt jedoch oft der von problemspezifischen Algorithmen hinterher. Sie bieten zudem keine Möglichkeit, um zu detektieren, ob ein Optimierungsproblem unzulässig ist, was die Wartung und Überprüfung der Korrektheit des Optimierungsmodells erschwert. Auch bieten sie keine Möglichkeit, um die Qualität der gefundenen Lösung im Vergleich zu der bestmöglichen Lösung abzuschätzen. Diese Lücke, also das verschenkte Potenzial, kann schon auf kleinen Instanzen signifikant sein.

2.2.2 EXAKTE VERFAHREN

Während sich im Feld der Heuristiken eine Vielzahl an verschiedenen Algorithmen entwickelt hat, hat sich bei den exakten Verfahren vor allem das **Branch-and-Bound-Verfahren** (Lawler und Wood 1966) durchgesetzt. Im Kern basiert das Verfahren wie die Brute-Force-Methode auf dem Enumerieren vieler Entscheidungen. Jedoch wird durch das geschickte Ausnutzen von sogenannten primalen und dualen Heuristiken die Menge an zu testenden Lösungen deutlich reduziert, sodass in immer mehr Anwendungen die geforderten Laufzeiten erreicht werden. Das Verfahren startet mit dem Lösen einer Vereinfachung des originalen Problems. Dabei werden gewisse Nebenbedingungen aufgeweicht, um daraus ein lineares Optimierungsproblem zu machen. Für diese Klasse gibt es exakte Optimierungsverfahren, wie das Simplex-Verfahren, welches meist in kurzer Zeit eine optimale Lösung für das vereinfachte Problem liefert. Diese Lösung ist in der Regel nicht zulässig für das Originalproblem, bietet aber eine Abschätzung für die Qualität der optimalen Lösung des Ausgangsproblems. Wenn nun mit Hilfe einer beliebigen primalen Heuristik eine zulässige Lösung generiert wurde, lässt sich durch den Wert dieser Lösung zusammen mit dem Wert des reduzierten Problems eine Abschätzung abgeben, wie gut die gefundene Lösung im Vergleich zu der bestmöglichen Lösung ist. Mit Hilfe dieser Abschätzung kann entschieden werden, ob es sich lohnt, weiter nach besseren Lösungen und Abschätzungen zu suchen. Falls die Abschätzung dem Wert der Lösung entspricht, ist mathematisch beweisbar eine globale Optimallösung für das Optimierungsproblem gefunden.

Falls die globale Optimallösung noch nicht gefunden wurde, startet das Verfahren mit dem Branching-Schritt. Dabei wird z. B. eine binäre Entscheidung in dem Problem ausgewählt und das Problem in zwei Teilprobleme zerlegt. In jedem der zwei Teilprobleme wurde ein anderer Ausgang der Entscheidung ausgewählt, dadurch entstehen zwei kleinere Optimierungsprobleme. Nun kann in beiden Teilproblemen wiederum neue, verbesserte Abschätzungen berechnet werden und mit Hilfe von Heuristiken nach neuen Lösungen gesucht werden. Dieses Verfahren kann nun sukzessiv

weitergeführt werden. Durch jeden Branching-Schritt entstehen zwei neue Teilprobleme und so entsteht ein Suchbaum, der sog. Branch-and-Bound-Baum, welcher exponentiell anwächst. Durch den Bounding-Schritt wird versucht genau das zu verhindern und Teilzweige des Baumes werden abgeschnitten, um den Suchraum zu verringern. Falls die Abschätzung eines Teilproblems schlechter ist als die beste gefundene Lösung, kann dieser Teil des Baums entfernt werden, weil dort auch in weiteren Branching-Schritten keine bessere Lösung gefunden werden kann.

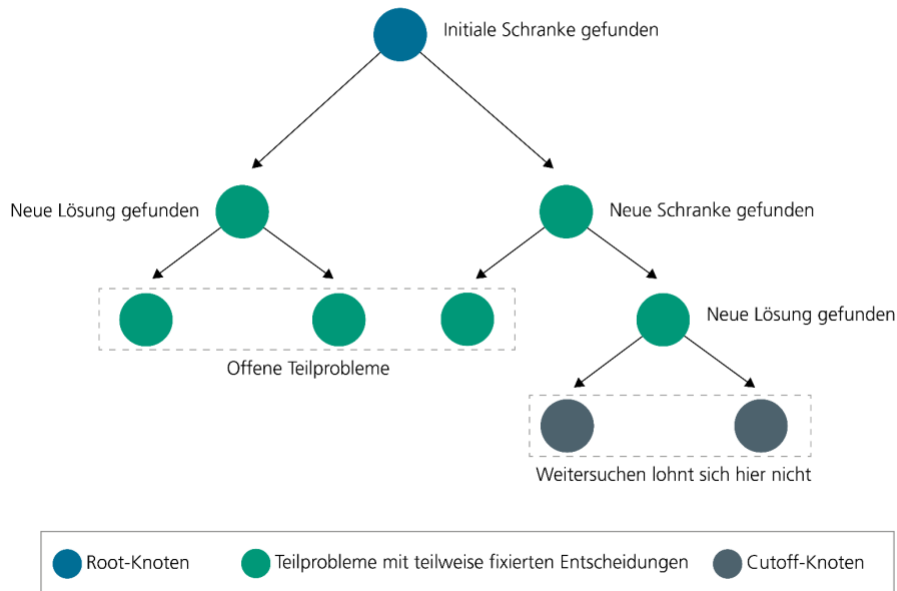


Abbildung 9: Graphische Darstellung eines Branch-and-Bound Baumverfahren

Der Branch-and-Bound Baum wird schematisch in Abbildung 9 dargestellt. Es werden sukzessive die Knoten des Baumes abgearbeitet und dort nach verbesserten Schranken und Lösungen gesucht. Abbildung 10 zeigt den Fortschritt des Verfahrens über die Zeit. Zu jedem Zeitpunkt zeigt die Differenz zwischen bester Schranke und Lösung das restliche Optimierungspotential. Mit fortschreitender Rechenzeit schließt sich diese Differenz, und das Verfahren konvergiert schließlich gegen eine optimale Lösung.

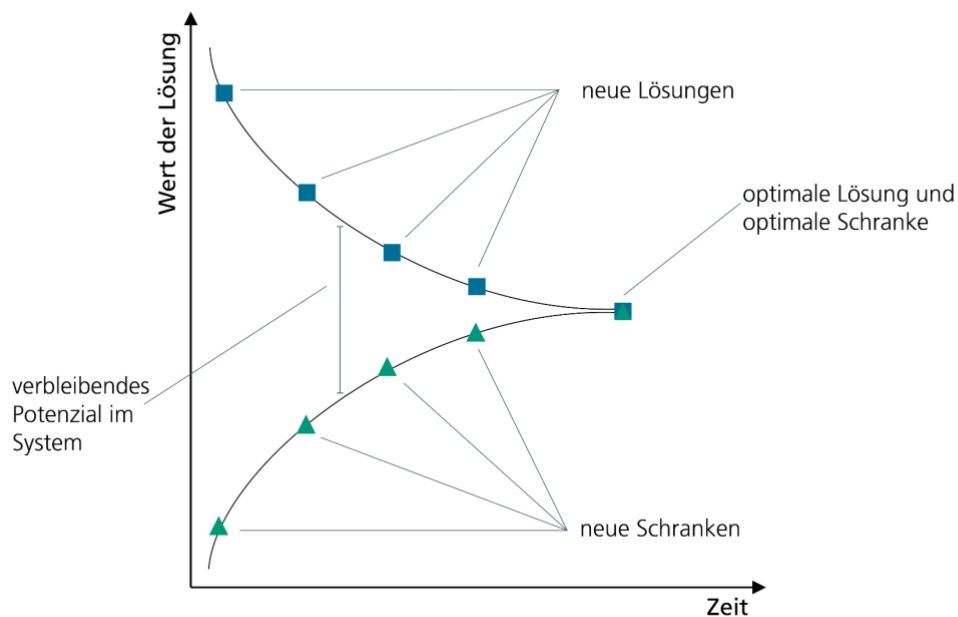


Abbildung 10: Graphische Darstellung eines Konvergenz-Optimierungsverfahrens

Es gibt viele verschiedene Verbesserungen des beschriebenen Basisverfahrens. So werden in der Praxis primale Heuristiken an den Knoten eingesetzt, die speziell für Branch-and-Bound entwickelt wurden und nach besseren zulässigen Lösungen zu suchen. Darüber hinaus gibt es die Erweiterung hin zu dem in der Praxis eingesetzten Branch-and-Cut Verfahren, in dem zusätzliche duale Heuristiken verwendet werden, welche versuchen bessere Abschätzungen zu liefern, um den Suchraum noch schneller zu reduzieren.

Die exakten Verfahren (insbesondere Verfahren vom Branch-and-Bound-Typ) bieten den Vorteil, dass sie zu jeder Zeit eine Abschätzung über die Qualität der gefundenen Lösung im Vergleich zur Optimallösung liefern. Dadurch kann jederzeit informiert entschieden werden, ob noch weitere Rechen- und Zeitkapazität in das Suchen einer besseren Lösung investiert werden sollte. Über die Laufzeit hinweg werden sowohl die Abschätzungen als auch die gefundenen Lösungen verbessert, sodass immer entschieden werden kann, ob sich weitere Rechen- und Zeitkapazität lohnen. Dadurch können insbesondere einfache Instanzen viel schneller gelöst werden, weil ein klares Abbruchkriterium besteht. Die Verfahren ermöglichen eine deutlich verbesserte Analysefähigkeit der Fragestellung, da mit Hilfe der Abschätzung das Potenzial im System für Weiterentwicklungen oder verschiedene Szenarien evaluiert werden kann. Die Verfahren bieten zudem die Möglichkeiten, auch Unzulässigkeiten im System zielgenau zu erkennen.

Während heuristische Verfahren auf Grund ihrer geringeren Einstiegshürden und freien Verfügbarkeit schon häufig eingesetzt werden, sind exakte Verfahren noch weniger verbreitet. Durch die enorme Weiterentwicklung der mathematischen Methoden in den letzten Jahrzehnten erschließen die exakten Verfahren jedoch immer mehr Einsatzgebiete. Gerade bei Anwendungen, die ein hohes wirtschaftliches Potenzial aufweisen, lohnt sich eine Evaluation und gegebenenfalls ein Umstieg auf exakte Verfahren auf Grund ihrer vielseitigeren Einsetzbarkeit und mächtigeren Analysefähigkeit. Praktische Fragestellungen führen oft zu komplizierten Optimierungsproblemen, für welche die Wahl des Optimierungsansatzes entscheidend ist. Die Entwicklung von problemspezifischen Algorithmen kann hier den Unterschied machen, ob die Optimierungsprobleme in der Praxis zufriedenstellend gelöst werden können und die entwickelten Softwaretools ihren maximalen Nutzen entfalten.

Zusammenfassend lässt sich sagen, dass Optimierung Unternehmen ermöglicht, einfachere, bessere und schnellere Entscheidungen zu treffen. Diese führen zu Kosteneinsparungen und verringern das wirtschaftliche Risiko. In jedem Fall lohnt es sich, Experten für mathematische Optimierung hinzuzuziehen, um eine fundierte Einschätzung für eine erfolgsversprechende Herangehensweise an die Fragestellung zu erhalten.

3 ERFOLGSGESCHICHTEN AUS DER PRAXIS

Wie die zuvor beschriebenen Methoden aus den Bereichen der mathematischen Optimierung und Data Science auf konkrete Fragestellungen im Bereich der Supply Chain angewandt werden können, soll im Nachfolgenden anhand von bereits erfolgreich abgeschlossenen und zu diesem Zeitpunkt noch laufenden Projekten der Arbeitsgruppe SCS vorgestellt werden. Die jeweils angegebenen Keywords sollen die Orientierung und die Suche nach relevanten Projekten erleichtern.

Tabelle 1: Überblick über Erfolgsgeschichten aus der Praxis. dient als Überblick der Data Analytics Projekte der Fraunhofer-Arbeitsgruppe eingeordnet entlang methodischer Forschungsbereiche.

	PRE- PROCESSING	CLUSTERING	KLASSISCHE STATIS- TISCHE VERFAHREN	NEURONALE NETZE	BAYES'SCHE VERFAHREN	HEURISTIKEN	EXAKTE VERFAHREN
KITE	X		X	X	X	X	X
PRODAB	X		X	X	X		
ALONE	X	X	X		X		X
MARTIN BAUER	X					X	X
MAGNA	X		X				
BSH	X	X	X	X	X		X
SCHNELL- ECKE	X					X	X

Tabelle 1: Überblick über Erfolgsgeschichten aus der Praxis.

3.1 LAUFENDE PROJEKTE

3.1.1 KITE: KÜNSTLICHE INTELLIGENZ IM TRANSPORT ZUR EMISSIONSREDUKTION

Keywords: Tourenplanung, Demand Forecasting, Emissionsreduktion, Green Logistics, Zeitreihenprognose, Straßengüterverkehr, Transportlogistik

Einleitung und Beschreibung der Fragestellung

Eine der großen Herausforderungen im Umgang mit dem Klimawandel ist die Senkung der Treibhausgasemissionen im Verkehr. Gerade der gewerbliche Güterverkehr hat dabei ein hohes Potenzial, die Emissionen zu senken. Denn: Ein beträchtlicher Anteil der LKW-Fahrten ist nicht optimal ausgelastet. Damit Disponenten Touren optimal auslasten können, akquirieren sie teilweise gezielt neue Aufträge oder verschieben Kundenaufträge. Darüber hinaus wurden in den letzten Jahren auch Frachtbörsen immer populärer. Alle diese steuernden Eingriffe erfordern eine Vorhersage, welches Frachtvolumen an den verschiedenen Stellen im Netzwerk anfallen wird, damit genug Zeit für die Absprache mit dem Kunden bleibt. Diese Prognose muss wiederum in die Tourenplanung integriert werden. Im Moment ist die Genauigkeit dieser Prognose vor allem von der Erfahrung des Disponenten abhängig. Ein datengetriebenes Verfahren kann hingegen für eine gleichbleibend hohe Prognosegenauigkeit sorgen und Disponenten die Zeit geben, sich um die echten Problemfälle zu kümmern.

Während schon viele Unternehmen Lösungen zur Tourenplanung anbieten, gibt es bisher keine zufriedenstellende Lösung, datengetriebene Vorhersagen in diese zu integrieren. Darüber hinaus liegt in der Thematik der Bedarfsprognosen von Frachtvolumen ein hohes Potential mit vielen offenen Forschungsfragen.

Ergebnisse des Projektes und Zusammenarbeit

Im Rahmen des Forschungsprojektes **KITE**, welches vom BMVI gefördert wird, entwickeln die Forscher*innen bis zum Jahr 2023 einen Algorithmus, welcher Frachtvolumina auf verschiedenen Ebenen und Horizonten vorhersagt. KITE schließt dabei an die Ergebnisse des Projekts **KIVAS** an, welches den Mehrwert externer Variablen auf die Prognose von Frachtvolumen untersucht hat. Im Rahmen von KITE kommen verschiedene Ansätze zur Zeitreihenvorhersage zum Einsatz. Insbesondere möchten sie eine Quantifizierung der Unsicherheit der Prognose ermitteln. Die Ergebnisse dieser Vorhersage sollen dann in das Tourenplanungsproblem integriert werden. Dabei berücksichtigen sie auch spezifische Restriktionen, wie die Lenkzeiten von Fahrern oder die Zuordnung von Fahrzeugen an transportierte Materialien. Neben der Optimierung von Touren kann die Prognose außerdem auch im strategischen Bereich genutzt werden, um gezielt die Struktur des Netzwerks des jeweiligen Transportunternehmens zu verbessern.

Das Projektteam hat das Ziel, den Anteil der Leerfahrten bei den betrachteten Betrieben, um bis zu 15 % zu reduzieren. An der Entwicklung und Pilotierung des Verfahrens sind neben dem Entwicklungspartner Optitool GmbH, die BLG Logistics Group AG & Co. KG sowie die Schmahl & Stoepel GmbH als Anwendungspartner an Bord.

Ausblick

Wenn sich im Forschungsprojekt der Mehrwert des Verfahrens zeigen lässt, hoffen die Forscher*innen, dass sich das entwickelte Verfahren in der Praxis durchsetzt und effektiv dazu beiträgt Emissionen zu reduzieren.

Weitere Informationen

Weitere Informationen zu diesem Projekt finden Sie unter:

<https://www.scs.fraunhofer.de/de/referenzen/kite.html>

3.1.2 PRODAB: DURCHLAUFZEITPROGNOSE MIT BAYES'SCHEN NETZEN

Keywords: *Process Mining, Durchlaufzeitprognose, Qualitätsmanagement, Intralogistik, Transportlogistik, resiliente Lieferketten, Bayes*

Einleitung und Beschreibung der Fragestellung

Eng getaktete Lieferketten erweisen sich im Störfall oft als fragil. Im schlimmsten Fall kommt es zu Bandstillständen, weil Zulieferprodukte nicht rechtzeitig geliefert werden können. Sendungen, welche Gefahr laufen nicht rechtzeitig geliefert zu werden, sollten deshalb möglichst früh im Prozess identifiziert werden. Die manuellen Nachverfolgung von hunderttausenden von Sendungen täglich ist allerdings mit unverhältnismäßig hohen Kosten verbunden. Die Wissenschaftler*innen entwickeln deshalb eine Software, mit der sich jede einzelne Sendung automatisiert mithilfe von Event-Logs (bspw. Scannungen) verfolgen lässt und die live die jeweilige verbleibende Durchlaufzeit prognostiziert. So können riskante Sendungen bereits früh im Prozess erkannt und entsprechende Gegenmaßnahmen eingeleitet werden.

Ergebnisse des Projekts und Zusammenarbeit

Die Durchlaufzeit von Einheiten in Prozessen ist von vielen verschiedenen Faktoren abhängig. Insbesondere für Prozesse, welche nicht am Fließband ablaufen, ist sie deshalb für einzelne Einheiten nicht trivial berechenbar. Gleichzeitig sind über Einheiten in Prozessen schon heute viele beschreibende Daten vorhanden: Regelmäßige Scannungen geben zum Beispiel Aufschluss über den aktuellen Prozessschritt und aus Auftragsdaten lässt sich der assoziierte Arbeitsaufwand berechnen. Darüber hinaus sind häufig Daten über auf den Prozess wirkende Größen wie beispielsweise das verfügbare Personal oder das aktuelle Auftragsbacklog vorhanden.

Um all diese – teilweise unternehmensindividuellen – Faktoren zu berücksichtigen verwenden die Forscher*innen Bayes'sche Netze, wie in Kapitel 2.1.1 „überwachtes Lernen“ erklärt. Die Struktur dieser Netze kann an den jeweiligen spezifischen Prozess angepasst werden. So lässt sich domänenspezifisches Wissen, wie beispielsweise Kausalitäten zwischen den Einflussgrößen, explizit modellieren. Die konkrete quantitative Ausprägung dieser Wirkzusammenhänge lernt das Verfahren dann aus vergangenen Sendungsdaten. Wenn beispielsweise ein höheres Auftragsbacklog zu einer höheren Durchlaufzeit führt, kann das Bayes'sche Netz diesen exakten Zusammenhang aus vergangenen Prozessdurchläufen lernen. Wenn das komplexe Zusammenspiel dieser verschiedenen Faktoren dazu führt, dass die Durchlaufzeit zu lange wird, also der gewünschte Liefertermin nicht gehalten wird, soll die PRODAB Software eine entsprechende Warnung ausgeben.

Im Rahmen des Projekts werden zwei verschiedene Anwendungsfälle betrachtet: Die Kommissionierung bei einem Maschinenbauunternehmen und der Umschlagsprozess bei einem Transportlogistikunternehmen. Für jeden der Use-Cases haben die Forscher*innen spezifische Bayes'sche Netze entwickelt. Parallel entwickelt das Team eine Software, mit der diese Bayes'schen Netze in die Unternehmens-IT-Infrastruktur eingebunden werden.

Ausblick

Wenn es gelingt, Wirkzusammenhänge in Prozessen akkurat mithilfe von Bayes'schen Netzen darzustellen, stellt sich im nächsten Schritt die Frage: Wie können diese Prozesse optimal gesteuert werden? Dazu ist wiederum die Kopplung von Vorhersage und mathematischer Optimierung notwendig. In einem Folgeprojekt planen die Forscher*innen dieses Thema erforschen zu dürfen.

Weitere Informationen

Weitere Informationen zu diesem Projekt finden Sie unter:

<https://www.scs.fraunhofer.de/de/referenzen/prodab.html>

3.1.3 ALONE: PRÄSKRIPTIVE ANALYSE DIGITALISierter WERTSCHÖPFUNGSKETTEN

Keywords: *Präskriptive Analyse, Gemischt-ganzzahlige nicht-lineare Optimierung, Multikriterielle Optimierung, Demand Forecasting*

Einleitung und Beschreibung der Fragestellung

Moderne Unternehmen agieren heutzutage in immer enger verwobenen Netzwerken. Einer der Haupttreiber dafür ist die Digitalisierung, die Unternehmen nach innen und auch über Unternehmensgrenzen hinweg miteinander verbindet und so hochgradig komplexe Systeme entstehen lässt. Gleichzeitig werden durch die digitalen Prozesse beständig Daten erzeugt, etwa in ERP- und MES-Systemen sowie durch die sensorische Erfassung an der Produktionsstätte. Im Rahmen

des ADA Lovelace Centers forscht die Arbeitsgruppe SCS in der Applikation ALONE daran, welche Arten an Fragestellungen durch die Verwendung solcher Datensätze beantwortet werden können. Unser Fokus liegt dabei auf der Planungs- und Entscheidungsunterstützung für die Manager*innen und Fachplaner*innen mittels sog. präskriptiver Analysen. Diese Unterstützung wird durch eine Kombination von Prognose und Optimierung besonders wirkmächtig, denn sie erlaubt es dem Unternehmen, auf die erwartete zukünftige Entwicklung, z. B. bei Preisen und Nachfrage, stets adäquat zu reagieren. Dazu gehen die Forscher*innen des Fraunhofer SCS und der Friedrich-Alexander-Universität Erlangen-Nürnberg so vor, dass sie die in der Industrie vorkommenden Problemstellungen so weit abstrahieren, dass sie deren mathematische Eigenschaften analysieren und geeignete Lösungsmethoden dafür entwickeln können. Dadurch können für die erzielten Lösungen mathematisch beweisbare Gütegarantien gegeben werden, man erhält hochgradig performante Berechnungsverfahren und es wird sichergestellt, dass die verwendeten Ansätze auch auf verwandte Fragestellungen übertragbar und somit allgemein einsetzbar sind.

Bei vielen Fragestellungen in Unternehmen können Prognose- und Optimierungsmodelle entlang der Wertschöpfungskette eine wertvolle Entscheidungsunterstützung liefern. In dieser Applikation werden daher Prototypenwendungen für industrierelevante Planungsaufgaben erforscht. Unser Ziel ist es, die sich daraus ergebenden Modelle mathematisch zu erforschen und geeignete Lösungsverfahren dafür zu entwickeln. In synergetisch ergänzenden, weiterführenden Projekten mit verschiedenen Bereichen der Industrie werden diese Prototypenwendungen dann in die Praxis überführt.

Herangehensweise und Ergebnisse des Projekts

Als Ausgangspunkt sei eine modellhafte Wertschöpfungskette innerhalb eines produzierenden Unternehmens in der verarbeitenden Industrie gegeben. Dabei setzen sich die Forscher*innen das Ziel, die Geschäftsindikatoren des Unternehmens zu optimieren, bspw. den Umsatz zu maximieren, die Kosten zu senken oder das Geschäftsrisiko zu minimieren, indem die zur Verfügung stehenden Stellschrauben der Produktion bestmöglich benutzt werden. In diesem Zusammenhang lassen sich mehrere Anwendungsfälle für Prognose- und Optimierungsmethoden betrachten, die zunächst als einzelne Module gelöst und anschließend miteinander verknüpft werden können. Beispiele hierfür sind die Bereiche Personaleinsatzplanung, Auslieferungslogistik oder Produktionsplanung. Am Ende sollen alle diese einzelnen Planungsschritte in einem ganzheitlichen Prozess optimiert werden.

In der Praxis werden die Pläne für die vorher genannten Fragestellungen häufig noch manuell erstellt und nicht mit Hilfe von mathematischer Software. Dies führt häufig zu zeitaufwändigen Planungsprozessen und zu suboptimalen Plänen. Die dahinterliegenden mathematischen Fragestellungen lassen sich sehr gut mit Hilfe der gemischt-ganzzahligen Optimierung (engl. mixed-integer programming, MIP) modellieren. Der direkte Einsatz dieser Modelle führt jedoch bei komplexen Problemen oft nicht zu der in der Praxis gewünschten Rechenzeitperformanz. In dem Projekt werden deswegen neue mathematische Ansätze erforscht, die zur effizienten Lösung dieser Entscheidungsprobleme beitragen.

Exemplarisch werden in diesem Projekt sogenannte Pooling- oder Mischungsprobleme betrachtet. Mischungsprobleme treten in der Produktionsplanung vieler fertigenden Unternehmen auf, und es müssen dabei Rohstoffe mit schwankende Qualitätseigenschaften auf optimale Weise zu

Endprodukten verarbeitet werden, deren Qualitätsparameter dann innerhalb gewisser definierter Schranken liegen müssen. Dabei ist insbesondere eine vorausschauende Lagerhaltung von großer Wichtigkeit. Auf diesem Gebiet werden deshalb neuartige, industriell relevante Problemvarianten erforscht und dafür effiziente Optimierungsverfahren entwickelt. Dazu analysieren die Forscherinnen und Forscher des Fraunhofer SCS und der Friedrich-Alexander-Universität Erlangen-Nürnberg die mathematische Struktur dieser Problemklasse. Die entwickelten Verbesserungen kommen zum Beispiel beim Teeproduzenten Martin Bauer GmbH (siehe Abschnitt 17) zum Einsatz, sind dabei jedoch allgemein anwendbar für die ganze Problemklasse und lassen sich einfach auf ähnliche Planungsprobleme in anderen Branchen adaptieren, in denen Mischungsfragestellungen auftreten.

Parallel dazu wird an dem generalisierbaren Einsatz der korrespondierenden Prognosefragestellungen geforscht. Typischerweise sind KI-Anwendungen auf einen speziellen Anwendungsfall zugeschnitten und auf Basis sehr konkreter Datensätze trainiert. Dadurch sind die verwendeten KI-Methoden in der Praxis nur schwer auf neue Anwendungsfälle zu übertragen und es können Schwierigkeiten bei der Inklusion neuer Datensätze auftreten. Ein möglicher Weg, mit dieser Problematik umzugehen und Algorithmen stabiler anwendbar zu machen, sind Ansätze des automatischen maschinellen Lernens (AutoML). Deshalb wird untersucht, wie Ansätze von AutoML für die Anwendungsklasse zur Modellierung von Maschinenausfallverhalten („Predictive Maintenance“) genutzt werden können. Hierzu wenden die Wissenschaftler*innen Ansätze von AutoML auf typische Methoden für den Anwendungsfall Predictive Maintenance an, d.h. Survival-, Klassifikations- und Regressionsmethoden und erreichen damit, dass die Hyperparameter der Methoden nicht für jeden Datensatz erneut manuell gefunden werden müssen.

Ausblick

Im Fortlauf des Projektes untersuchen wir die dynamische Erweiterung des Optimierungsmodells, bei der die Zeit, die zwischen den Produktionsschritten vergeht, mit modelliert wird. Diese Erweiterung erlaubt es, in Echtzeit auf neue Informationen (bspw. sich ändernde Aufträge oder neue Laborergebnisse der Rohmaterialqualitäten) zu reagieren. Hierzu ist es zunächst notwendig algorithmische Verfahren herzuleiten, um echtzeitfähige Rechenzeiten zu erlauben. Außerdem soll die Nachfrage nach den Endprodukten z. B. mittels Bayesianischer Modelle prognostiziert werden. Diese Prognosen können für die Beschaffung von Rohstoffen genutzt werden, die wiederum einen Einfluss auf die Lösung des Produktionsproblems haben, da so durch die Optimierung der richtige Lagerbestand erreicht werden kann. In diesem Modus sollen nach und nach einzelne Elemente der Wertschöpfungskette modelliert und miteinander verknüpft werden um den Wertschöpfungsprozess immer weiter datengetrieben zu optimieren.

Weitere Informationen

Weitere Informationen zu diesem Projekt finden Sie unter:

<https://www.scs.fraunhofer.de/de/referenzen/ada-center/logistische-oekosysteme.html>

3.1.4 RESSOURCENSCHONENDE UND EFFIZIENTE PRODUKTIONSPLANUNG BEIM TEEZULIEFERER MARTIN BAUER

Keywords: *Pooling-Problem, Teemischung, Blending Problem, Mischproblem, Gemischt-ganzzahlige nicht-lineare Optimierung, Multikriterielle Optimierung, Lebensmittel, Qualitätsmanagement*

Einleitung und Beschreibung der Fragestellung

Die Martin Bauer Group befasst sich schwerpunktmäßig mit der Veredelung von pflanzlichen Rohstoffen zu Lebensmittel- und Arzneimitteltees, Extrakten, pflanzlichen Pulvern und Rohstoffen für die Getränke-Industrie. Die Produkte der Martin Bauer Group finden jedoch auch in diversen weiteren Lebensmittel-Industrien Anwendung. Die Rohwaren dafür werden global beschafft. Die Kunden sind in diesem Fall Handelsmarken, welche die fertigen Kräuter- und Früchtetee-mischungen vermarkten und vertreiben. Dabei stellt das Unternehmen sowohl Tees für den Lebensmittelbereich her, welche dann z. B. im Supermarkt verkauft werden, als auch Arzneimitteltees, z. B. zum Verkauf in Apotheken.

Der Fokus dieses gemeinsamen Projektes zwischen der Arbeitsgruppe SCS, der FAU Erlangen-Nürnberg und der Martin Bauer Group liegt auf der Produktion für den Bereich Kräuter- und Früchtetee-Mischungen. Die fertigen Teemischungen müssen diverse Qualitätsanforderungen bezüglich ihrer Inhaltsstoffe erfüllen. Typischerweise müssen die relevanten Inhaltsstoffe innerhalb vorgegebener Konzentrationsbereiche liegen, entsprechend den jeweiligen Kunden –und Marktanforderungen. Die eingekauften pflanzlichen Rohstoffe weisen wie alle natürlichen Produkte Schwankungen in den Qualitätsparametern auf, da diese von vielen Faktoren beim Anbau abhängen, wie z. B. der Bodenbeschaffenheit dem örtlichen Klima oder auch den Wetterbedingungen eines gegebenen Jahres. Deshalb ist eine zeitaufwändige und genaue Planung des Einsatzes der verfügbaren Lieferungen an Rohware für die verschiedenen Produktionsaufträge notwendig, um die Kunden termingerecht gemäß ihren jeweiligen Anforderungen zu beliefern zu können. Um dies jedoch auch langfristig zu gewährleisten, ist es notwendig, die strategischen Auswirkungen der Planung zu betrachten. So muss beim Einsatz des verfügbaren Rohmaterials insbesondere auf eine vorausschauende Lagerführung geachtet werden.

Ergebnisse des Projektes und Zusammenarbeit

Die Martin Bauer Group stellt in diesem Projekt die notwendigen Daten sowie das Know-how über die Details der Produktionsplanung bereit. Das in der Fragestellung beschriebene Optimierungsproblem wird in der Mathematik als Pooling-Problem bezeichnet und tritt in vielen Anwendungen auf, bei denen natürliche Produkte gemischt werden müssen, wie z. B. der Erdölverarbeitung oder dem Abwassermanagement. Das Lösen von Pooling-Problemen ist auf Grund seiner Komplexität und seiner kombinatorischen Reichhaltigkeit ein sehr aktives Forschungsgebiet.

Die zu erstellende Optimierungssoftware erhält als Eingaben einerseits Informationen über die aktuellen Lagerbestände, wie z. B. Mengen, Lagerdauer und Laboranalytik der verfügbaren Rohmaterialien, sowie der Zwischen- und Endprodukte. Zusätzlich liest die Software Rezepturen der Zwischen- und Endprodukte ein, welche beschreiben, wie viel von jedem Rohmaterial oder Zwischenprodukt je produzierter Einheit eingesetzt werden muss. Weiterhin ist eine Liste an Kundenaufträge gegeben. Für jeden dieser Aufträge ist das bestellte Endprodukt, die davon zu produzierende Menge sowie die einzuhaltenden Qualitätsparameter hinterlegt.

Die Optimierungssoftware ist auf Basis dieser Daten in der Lage, verschiedene Optimierungsmodelle aufzustellen, welche optimale Produktionspläne nach vorab definierten Zielkriterien erstellen. Für eine gegebene Auftrags- und Lagersituation erhält der Benutzer als Ausgabe für jeden Auftrag die

Produktionsmenge und die erwarteten Inhaltsstoffwerte. Bei Aufträgen, welche der optimierte Plan nicht oder nur teilweise zu produzieren vorsieht, gibt die Software eine Erklärung für das erzielte Ergebnis. So könnte es z. B. sein, dass die aktuell verfügbaren Rohwaren nicht ausreichen oder bereits für andere Aufträge genutzt werden. Es könnte auch sein, dass die Rohwaren zwar ausreichen, aber dass damit nicht die geforderte Produktqualität erreicht werden kann. In letzterem Fall die Software Auswertungen darüber, welche Inhaltsstoff-Minimal- und Maximalwerte nicht eingehalten werden können und welche Werte erreicht werden im besten Fall erreicht werden können. Zusätzlich erhalten die Planer einen Bericht über den nach Umsetzung des Produktionsplans resultierenden Lagerbestand, um auch die langfristigen Auswirkungen des Plans in die Entscheidung miteinbeziehen zu können. Die Wählbarkeit verschiedener Zielkriterien ermöglicht es darüber hinaus, bei der Optimierung z. B. einen Fokus auf eine möglichst hohe Auftragsabdeckung zu legen, oder auf eine zügige Verwendung der bereits länger lagernden Rohwaren, um Qualitätsverluste zu minimieren.

Die Produktionsplaner der Martin Bauer Group sind mit Hilfe der Software in der Lage, schnell unterschiedliche Szenarien durchzuspielen und diese fundiert zu bewerten, wodurch diese in ihrer Produktivität unterstützt werden. Große Teile der zeitlich aufwändigen, kostenintensiven und fehleranfälligen Detailplanung fallen dadurch weg, und die Planer können sich auf weiterreichende, strategische Fragen der Planung konzentrieren. Vor allem aber besteht ein signifikantes Kosteneinsparpotenzial dadurch, dass dem Planer effiziente Handlungsoptionen vorgeschlagen werden, die auf Grund der kombinatorischen Vielfalt des unterliegenden Optimierungsproblems vom Menschen schwer zu finden sind.

Ausblick

Für die oben beschriebene Produktionsfragestellung gibt es vielfältige Möglichkeiten der Weiterentwicklung. Da der Produktionsplan jeden Tag aktualisiert werden muss und die vorangegangenen Entscheidungen selbstverständlich Auswirkungen auf Folgepläne haben, ergibt sich weiteres Optimierungspotenzial, durch das Ausnutzen von Synergien im zeitlichen Verlauf der Planung. Zudem kann auf Basis von Prognosen der zukünftigen Auftragsentwicklung bereits Produkte vorproduziert werden, um eine schnellere Verfügbarkeit zu gewährleisten. Auch besteht die Möglichkeit das aktuell untersuchte Problem um den Aspekt der Maschinenbelegung zu erweitern. Dadurch wird es zusätzlich möglich die Reihenfolge, in der die Aufträge auf den vorhandenen Maschinen produziert werden sollen, zu optimieren.

Weitere Informationen

Weitere Informationen finden Sie unter:

<https://www.scs.fraunhofer.de/de/referenzen/ada-center/logistische-oekosysteme.html>

<https://en.www.math.fau.de/edom/projects-edom/logistics-and-production/optimized-production-in-the-tea-industry/>

3.2 ABGESCHLOSSENE PROJEKTE

3.2.1 FEHLERQUELLEN ERKENNEN – FEHLER VERMEIDEN

Keywords: *Produktionsprozess, Produktionsfehler, Fehler-Ursachen-Analyse, Process Mining, Klassifikation, logistische Regression, Bayes, Entscheidungsbäume, Imbalanced Data.*

Einleitung und Beschreibung der Fragestellung

Magna International Inc. ist ein kanadischer Automobilzulieferer und betreibt mit seinen knapp 170.000 Beschäftigten mehr als 400 Entwicklungs- und Produktionszentren. Als Deutsches Tochterunternehmen für die Sitzfertigung betreibt die Magna Seating International (Germany) GmbH in Neuburg a. d. Donau ein Werk zur Montage von Autositzen. Diese werden bedarfsgetrieben mit geringem Vorlauf an einen Automobilhersteller in Ingolstadt geliefert. Die möglichst frühe Erkennung sowie gänzliche Vermeidung von Fehlern in der Montage ist eines der Hauptbestreben der Verantwortlichen im Werk. Damit soll u.a. die Termintreue eingehalten sowie eine hohe Produktqualität sichergestellt werden. Ziel des Projektes mit den Forschenden der Arbeitsgruppe SCS war es, mithilfe der von Magna erfassten Daten aus der Produktion Auffälligkeiten und mögliche Ursachen für Fehlerfälle zu identifizieren. Die gewonnenen Erkenntnisse können im Anschluss als Ausgangspunkt für Verbesserungen des Montageprozesses genutzt werden.

Ergebnisse des Projektes und Zusammenarbeit

Bevor das Aufkommen von Fehlerfällen analysiert werden konnte, musste zunächst der Datenbestand gesichtet, aufbereitet und zusammengeführt werden. Dies geschah in regem Austausch mit den IT- und Produktionsexperten von Magna, um deren umfangreiches Domänenwissen zu nutzen. Mit univariaten Auswertungen verschafften sich die Data-Science-Experten der Arbeitsgruppe SCS einen ersten Überblick über die Variablen im entstandenen Datensatz. Parallel wurde der Montageprozess aufgenommen und visualisiert, um mögliche Zusammenhänge besser zu verstehen. Hier wurde eine Darstellung, die durch das vorhandene Expertenwissen bei Magna entstand, einem aus den Daten via Process Mining erstellten Graphen (einer sog. Process Map) gegenübergestellt. Der Vergleich zwischen erwartetem und tatsächlichem Ablauf des Montageprozesses lieferte erste Erkenntnisse, welche Arbeitsschritte noch nicht präzise in den MES-Daten erfasst werden und wo die vorgesehene Abfolge von Montageschritten nicht eingehalten wird.

Nach Absprache mit Magna und unter Berücksichtigung der Darstellungen des Montageprozesses entschieden sich die Data-Science-Experten der Arbeitsgruppe SCS für Variablen, deren Einfluss sie auf das Vorkommen von Fehlerfällen untersuchten. Dafür wurden zunächst verschiedene univariate Analysen durchgeführt, also der Zusammenhang von jeweils einem Feature mit dem Auftreten von Fehlerfällen untersucht, um darauf basierend ein Klassifikationsmodell mit den stärksten Einflussvariablen zu erstellen.

Dieses Klassifikationsmodell sagt anhand der Merkmale eines Sitzes vorher, ob dieser ein Fehlerfall ist oder nicht. Es handelt sich also um eine binäre Klassifikation. Um ein solches Modell zu schätzen bzw. zu lernen wurden verschiedene Methoden aus der Statistik und ML verwendet und verglichen. Die erste verwendete Methode ist die klassische logistische Regression, die eine Variante der generalisierten linearen Modelle darstellt, bei der die Logit-Funktion als Linkfunktion gebraucht wird und dadurch die Wahrscheinlichkeiten der Klassenzugehörigkeit vorhersagt. Zusätzlich wurde dieses Modell auch auf Bayesianischem Wege geschätzt.

Außerdem wurde noch die Performanz verschiedener weiterführender Ansätze verglichen, die auf Entscheidungsbäumen basieren: eine Implementierung mit Gradient-Boosting, eine Weiterentwicklung davon speziell für die Behandlung kategorialer Variablen und ein sogenannter Isolation Forest, der für das Aufspüren von Anomalien entwickelt worden ist.

Eine entscheidende methodische Herausforderung dieses Projektes war die niedrige Anzahl an dokumentierten Fehlerfällen. Obwohl eine geringe Menge an Fehlerfällen erst einmal positiv zu betrachten ist, haben die Verfahren im Umkehrschluss leider auch wenig Möglichkeit, zu lernen, welche Kombination aus Merkmalen bei Sitzen mit Fehler vermehrt auftreten. Außerdem hat ein

Modell, das nie Fehlerfälle vorhersagt, eine hohe Trefferquote, obwohl es keinen Informationsgewinn liefert. Deswegen wurde mit jeder Methode jeweils dasjenige Modell gesucht, welches das höchste sogenannte F1-Maß hat. Dieses Maß berücksichtigt gleichermaßen den Anteil an Nicht-Fehlerfällen, die korrekt als solche erkannt werden, und den Anteil an Fehlerfällen, die korrekt als solche erkannt werden.

Die so geschätzten bzw. gelernten Modelle wurden im Anschluss miteinander und mit einem Benchmarkmodell, d.h. mit einem bestmöglichen Referenzmodell, verglichen. In diesem Fall wurde der Anteil an Fehlerfällen in den Daten als Wahrscheinlichkeit, dass ein Sitz einen Fehler hat, verwendet, ohne Einflussvariablen in Betracht zu ziehen. Zum Vergleich der Modelle wurde zum einen die Konfusionsmatrix verwendet, eine Vierfeldertafel, die die wahre Zuordnung zu Fehlerfall oder kein Fehlerfall in den Daten der prognostizierten Zuordnung gegenüberstellt. So kann gut überblickt werden, wie viele richtige und falsche Entscheidungen das jeweilige Modell getroffen hat. Zusätzlich wurde noch der Matthews-Correlations-Coeffizient (MCC) berechnet, der die Konfusionsmatrix auf eine Maßzahl zusammenfasst, die dann gut verglichen werden kann. In diesem Fall haben die beiden logistischen Regressionsmodelle und der Gradient-Boosting-Ansatz für kategoriale Variablen besonders gute Ergebnisse geliefert.

Ausblick

Aufgrund dieser Auswertungen war es für die Magna Seating möglich, den potenziellen Nutzen von ML-Ansätzen abzuschätzen. Es wurde identifiziert, dass vor allem automatisches maschinelles Lernen (AutoML) sowie die Erklärung der von KI getroffenen Entscheidungen zu einem Mehrwert führen können. So ist es durch die entwickelte Lösung künftig möglich, dass Magna über alle Produktionsstandorte hinweg einheitliche Analysen durchführt, die zu einer erhöhten Prozessstabilität und damit erhöhter Effizienz führen. Die umfassenden Datenanalysen und -aufbereitungen durch die Forscher*innen der Arbeitsgruppe SCS sind der nächste Schritt hin zu einer datengetriebenen Prozessvorhersage und -optimierung.

3.2.2 LANGFRISTIG ZUVERLÄSSIG: ERSATZTEILBEDARFE FÜR DAS NÄCHSTE JAHRZEHT VORHERSAGEN

Keywords: *Ersatzteil-Management, Demand Forecasting, End of Production, Schlusseindeckung, Obsoleszenz-Management, Clustering, Growth-Curve-Analysis*

Einleitung und Beschreibung der Fragestellung

Die BSH Hausgeräte GmbH ist eines der weltweit führenden Unternehmen der Branche und der größte Hausgerätehersteller in Europa. Weltweit beschäftigt die BSH knapp 60.000 Mitarbeiter*innen und unterhält ein weltweites Vertriebs- und Kundenetz, zu dem auch eine Reihe an Zentrallägern für Ersatzteile gehören. Wesentliches Element des Qualitätsversprechens von BSH Hausgeräte GmbH ist nämlich die lange Lebensdauer ihrer produzierten Geräte. Dazu gehört auch die Gewährleistung von Ersatzteilverfügbarkeit um Reparaturen und Wartungen langfristig durchführen zu können. Dies stellt das Unternehmen vor eine große Herausforderung: Bereits zum Ende der regulären Produktion müssen Disponent*innen abschätzen, wie viele Ersatzteile der Produkte in den nächsten Jahren nachgefragt werden.

Ergebnisse des Projektes und Zusammenarbeit

Disponent*innen entscheiden traditioneller Weise über die einzulagernden Ersatzteilmengen auf Basis einer detaillierten Kenntnis über die Artikel und damit verbundenen Erfahrungswerten. Dabei griffen sie auf ihre Erfahrungen aus bereits erfassten Produktlebenszyklen zurück und berücksichtigten die bisherigen Absätze des jeweiligen Ersatzteils. Dieses Verfahren ist jedoch zeitaufwändig, da alle relevanten Informationen aus verschiedenen Oberflächen und Systemen zusammengetragen werden müssen.

In Zusammenarbeit mit der Arbeitsgruppe SCS erkannten die Forscher*innen schnell, dass die Aufgabenstellung aus verschiedenen Gründen auch datengetrieben nicht trivial lösbar ist: Stammdaten mussten durch geschickte Verknüpfung und Heuristiken aufgefüllt werden. Zudem waren die klassischen Methoden zur Zeitreihenvorhersage nicht anwendbar. Die Disponent*innen unterstützten deshalb zunächst die KI-Experten mit ihrer Fachkenntnis, sodass gemeinsam eine „saubere“ Datenbasis geschaffen werden konnte. Aufbauend auf dieser Datenbasis entwickelten letztere ein neues Prognoseverfahren für Ersatzteile, indem sie aus verschiedenen Clustern von ähnlichen Ersatzteilen typische Absatzkurven ableiteten. Wenn nun der Bedarf bzw. die Nachfrage nach einem neuen Ersatzteil vorhergesagt werden soll, wird dieses auf Basis seiner Stammdaten einem bestimmten Lebenszyklus-Cluster zugeordnet. Dieses Clustering wurde iterativ wiederholt, sodass die datengetriebene Prognose inzwischen eine wertvolle Hilfe für die Disponent*innen darstellt. Damit das hier entwickelte Prognoseverfahren einfach anzuwenden ist, entwickelten die Forscher*innen eine graphische Oberfläche, die neben der Prognose auch die wichtigsten entscheidungsrelevanten Daten zu einem bestimmten Ersatzteil anzeigt. Über dieses Prognose-Tool können sich Disponent*innen zusätzlich Berichte zu allen Ersatzteilen generieren lassen. So kann beispielsweise auch Verschrottungspotenzial erkannt werden, indem das System Artikel identifiziert, denen voraussichtlich kein Bedarf mehr gegenübersteht. In den letzten beiden Jahren hat sich dieses Tool bewährt und alle Beteiligten arbeiten an seiner Erweiterung – die letzte Entscheidung über den Ersatzteilbestand verbleibt allerdings immer bei den Disponent*innen.

Ausblick

Neben der reinen Prognose von Ersatzteilbedarfen ist für BSH die tatsächlich einzulagernde Menge interessant. Wie der Fachmann weiß, ist die Lagerhaltung einem Zielkonflikt unterworfen: Werden zu viele Produkte eingelagert, entstehen unnötige Lagerkosten oder es müssen sogar Produkte verschrottet werden. Werden zu wenig Produkte eingelagert, können Kundenbedarfe nicht befriedigt werden, was gegebenenfalls zu hohen Opportunitätskosten führt. Gegeben einer Prognose können Verfahren der mathematischen Optimierung die wirtschaftlichste Entscheidung in diesem Zielkonflikt ermitteln. Idealerweise geht das wiederum einher mit einer Abschätzung der Unsicherheit in der Prognose. In einem Folgeprojekt setzt die Arbeitsgruppe SCS diese mathematische Optimierung für BSH um.

Weitere Informationen

Aus dem Projekt ist eine Veröffentlichung hervorgegangen, welche tiefere Einblicke zur Methodik geben kann: Menden, C., Mehringer, J., Martin, A. & Amberg, M., (2019). Vorhersage von Ersatzteilbedarfen mit Hilfe von Clusteringverfahren. HMD Praxis der Wirtschaftsinformatik: Vol. 56, No. 5. Springer. (S. 1000-1016). DOI: 10.1365/s40702-019-00532-7

3.2.3 DYNAMISCHE LAGERPLANUNG BEI SCHNELLECKE

Keywords: Lagermanagement, Intralogistik, Kommissionierung, Gemischt-ganzzahlige Optimierung, MIPS, Divide-and-Conquer, Dekomposition

Einleitung und Beschreibung der Fragestellung

Für Logistikdienstleister und Zulieferer ist der Umschlag von Waren an viele, teils stark ineinandergreifende Nebenbedingungen geknüpft, deren Nichteinhaltung zu kleineren Stopps im Betriebsablauf aber auch zu größeren Katastrophen führen kann. Dabei stellt sich nicht nur die Frage, welche Waren in welche Regale passen, es geht vor allem auch um Sicherheitsvorschriften für die Mitarbeiter und das Einhalten von Last- und Brandschutzbedingungen.

In der Regel setzt ein festgelegtes System die Einhaltung dieser Restriktionen um, welches allerdings viele Parameter wie z. B. die Entfernung zur Be- und Entladezone nicht berücksichtigt. Es entstehen Herausforderungen, derer sich die Firma Schnellecke Group AG & Co. KG gemeinsam mit der Fraunhofer-Arbeitsgruppe SCS des Fraunhofer IIS angenommen hat. Durch die Kombination von Logistik-Expertise auf der einen Seite und Kompetenzen der mathematischen Optimierung auf der anderen Seite wurde eine dynamische Lagerplanung erarbeitet.

Schnellecke ist ein Logistikdienstleister mit Sitz in Wolfsburg. Das Familienunternehmen bietet vor allem Services für die Automobilindustrie und ist weltweit an mehr als 70 Standorten mit einer Hallenfläche von 1.500.000 Quadratmetern vertreten. Die Unternehmensgruppe beschäftigt rund 20.000 Mitarbeiter.

Am Pilot-Standort Leipzig von Schnellecke erreichen täglich rund 1.000 bis 1.400 Klein- und Großladungsträger den Wareneingang bei rund 7.800 Handling Units im Hauptlager. Die bisherigen Systeme berücksichtigen zwar bestimmte Restriktionen, lassen viele Parameter aber außer Acht, die den späteren Betriebsablauf verbessern können. Ziel des Projektes war also eine optimale Verteilung von Klein- und Großladungsträgern im Hauptlager.

Ergebnisse des Projektes und Zusammenarbeit

Zuerst wurde von der Arbeitsgruppe SCS gemeinsam mit Schnellecke ein grundsätzliches Optimierungsmodell sowie die notwendigen Parameter und Rahmenbedingungen erarbeitet. Diese Informationen wurden von den Experten der Arbeitsgruppe in ein erstes Modell überführt, das im Anschluss mit tatsächlichen Industriedaten gefüttert wurde und dessen Ergebnisse auf ihre Validität hin geprüft wurden. Die Logistik-Expertise und jahrelangen praktischen Erfahrungen von Schnellecke halfen, um das Optimierungsmodell hinsichtlich realitätsnaher Abläufe im Produktionsalltag zu verfeinern.

Es handelt sich bei dem angewandten Verfahren um ein sogenanntes gemischt-ganzzahliges Optimierungsmodell, das mittels einer Methode gelöst wird, die das Optimierungsproblem sukzessive in immer kleinere Teilprobleme zerlegt, die dann mit einfacheren Verfahren gelöst werden können. Um eine Lösung für das eigentliche Problem zu finden, müssen die gefundenen Teillösungen nun wieder in umgekehrter Reihenfolge zusammengesetzt werden. Wie fein die Aufgabe unterteilt wird, hängt davon ab, wie viel Zeit man dem Algorithmus gibt – je mehr, desto feiner, desto besser die Lösung. Der Vorteil dabei ist außerdem, dass zu jedem Zeitpunkt des Optimierungsverlaufs bekannt ist, wie viel Verbesserungspotenzial im besten Fall noch erreicht werden kann.

Der Algorithmus zielt also zum einen darauf ab, sämtliche angelieferten Waren ordnungsgemäß im Lager zu verstauen, und zum anderen sollen die Prozesszeiten für die Lagerbewirtschaftung minimiert werden, um so einen effizienteren Arbeitsprozess zu erreichen. Dabei berücksichtigt der Algorithmus nicht nur sämtliche Einschränkungen, die bei der Warenlieferung und -platzierung beachtet werden müssen, er versucht auch die Zugänglichkeit zu den Waren optimal zu berechnen.

All das berechnet das Programm nahezu in Echtzeit und hilft damit, die Lagerführung sicherer und effizienter zu gestalten.

Ausblick

Die Ergebnisse der Forschung sollen nun in die Praxis überführt werden. In einem Anschlussprojekt sollen die errechneten Zeiteinsparungen mit der Anwendung in einem Lager der Schnellecke Logistics bestätigt werden. Dabei ergeben sich wiederum neue Aufgaben- und Fragestellungen: beispielsweise die Software-Integration in die bestehenden IT-Systeme oder wenn es durch unvorhergesehene Ereignisse zu einer Differenz zwischen Datenbasis des Modells und der tatsächlichen Situation im Lager kommt. Wenn auch diese Herausforderungen gelöst sind, steht einer Ausweitung auf weitere und größere Schnellecke-Standorte nichts im Wege.

Weitere Informationen

Weitere Informationen zu diesem Projekt finden Sie unter:

<https://www.scs.fraunhofer.de/de/referenzen/dynla.html>

4 WEITERE / KOMMENDE EXPERTISE IM BEREICH SUPPLY CHAIN ANALYTICS

Der Anspruch der Forschungsgruppen „Optimization“ und „Data Science“ der Fraunhofer-Arbeitsgruppe für Supply Chain Services ist, die angewandte Forschung im Bereich Supply Chain Analytics voranzutreiben und so die Forschungsgrenze immer weiter zu verschieben. Dafür fokussieren sie sich unter anderem auf die nachfolgend vorgestellten methodischen Themengebiete.

4.1 INTEGRATION VON DOMÄNENWISSEN IN FORECASTING

Im Gegensatz zum Modebegriff „Big Data“, liegt die Herausforderung im Bereich Forecasting in der Regel weniger darin, eine *Vielzahl* von Daten nutzbringend auszuwerten, sondern darin, dass für einen Sachverhalt nur wenig Beobachtungen vorhanden sein können – also aus *wenigen* Datenpunkten Mehrwert generiert werden soll. Gleichzeitig besitzen Experten und Facharbeitende jedoch häufig Domänenwissen über die zu prognostizierenden Sachverhalte. Mit diesem Domänenwissen können Menschen einerseits die Prognosegenauigkeit von Algorithmen verbessern. Zum anderen können sie auch erkennen und eingreifen, wenn die computergestützten Systeme einen Fehler machen – sofern die Interpretierbarkeit der Modelle gegeben ist. Ein manueller Eingriff in Modelle bedeutet immer eine Verschiebung der erlernten Modellparameter, welche die Modelle automatisiert annehmen sollten. Dieser Prozess wird auch als Domain Shift Adaption bezeichnet.

Bei manchen Modellen können Menschen mit Domänenwissen Strukturen und Abhängigkeiten vorgeben, sodass nur noch deren quantitative Ausprägung anhand von Daten geschätzt werden muss. Deshalb sind auch weniger Daten vonnöten, um diese Modelle zu schätzen. Ein weiterer Vorteil dieses Ansatzes ist, dass die so aufgestellten Modelle auch von Menschen ohne Statistikenkenntnisse gut interpretiert werden können. Insbesondere erforschen die Wissenschaftler*innen der Arbeitsgruppe SCS hier den Einsatz von Bayes-Netzen und Neuronalen Netzen.

Beispielsweise kann der Einfluss von Variablen aufeinander mit gewissen Annahmen vorgelegt werden: So könnte ein Eisverkäufer davon ausgehen, dass er mehr Eis verkauft je wärmer es wird – jedoch nur bis zu einer bestimmten Grenze. Bei Bayesianischen Modellen kann darüber hinaus eine Vermutung über die quantitative Ausprägung der Modellparameter integriert werden. Beispielsweise wissen Disponenten häufig, welche Bedarfszahlen für ein Produkt plausibel sind und welche sehr ungewöhnlich wären. Dieses Wissen kann dann als Ausgangspunkt für einen Mittelwert oder eine Varianz im Modell genutzt werden.

Schon bei der Wahl, welche Einflussvariablen bzw. Features überhaupt für eine Fragestellung berücksichtigt werden sollten, kann Domänenwissen einen wichtigen Beitrag leisten. Hier ist besonders Gegenstand der Forschung, wie sich datengetriebene Ansätze zur Feature Selection mit Domänenwissen kombinieren lassen, damit sie sich im Idealfall gegenseitig ergänzen und korrigieren.

Auch das Wissen über Ähnlichkeiten oder Hierarchien zwischen Produkten oder Waren kann sehr wertvoll für eine Analyse sein. Ist den Experten z. B. bekannt, dass mehrere Produkte ein ähnliches Absatzverhalten haben, können die Daten und damit die vorhandenen Informationen über diese Produkte gebündelt werden. So können Prognosen auch für Produkte verlässlich sein, über die selbst nur wenige Daten vorliegen oder die vielleicht vollkommen neu sind. Auch hierfür bieten sich

Bayesianische Ansätze an, weil hier Parameter geteilt oder vererbt werden können (Elias Gyftodimos und Peter A. Flach 2002, siehe auch Abbildung 11). Lange Zeit war der hohe Rechenaufwand bei Bayesianischen Methoden ein großes Hindernis. Aktuell wird erforscht, inwiefern verbesserte Optimierungsverfahren und höhere Rechenleistungen diese Methoden auch bei großen Datenmengen erfolgreich einsetzbar machen. Eine besonders interessante Fragestellung ist dabei die Prognose von Nachfrage im Einzelhandel bei großen, internationalen Unternehmen.

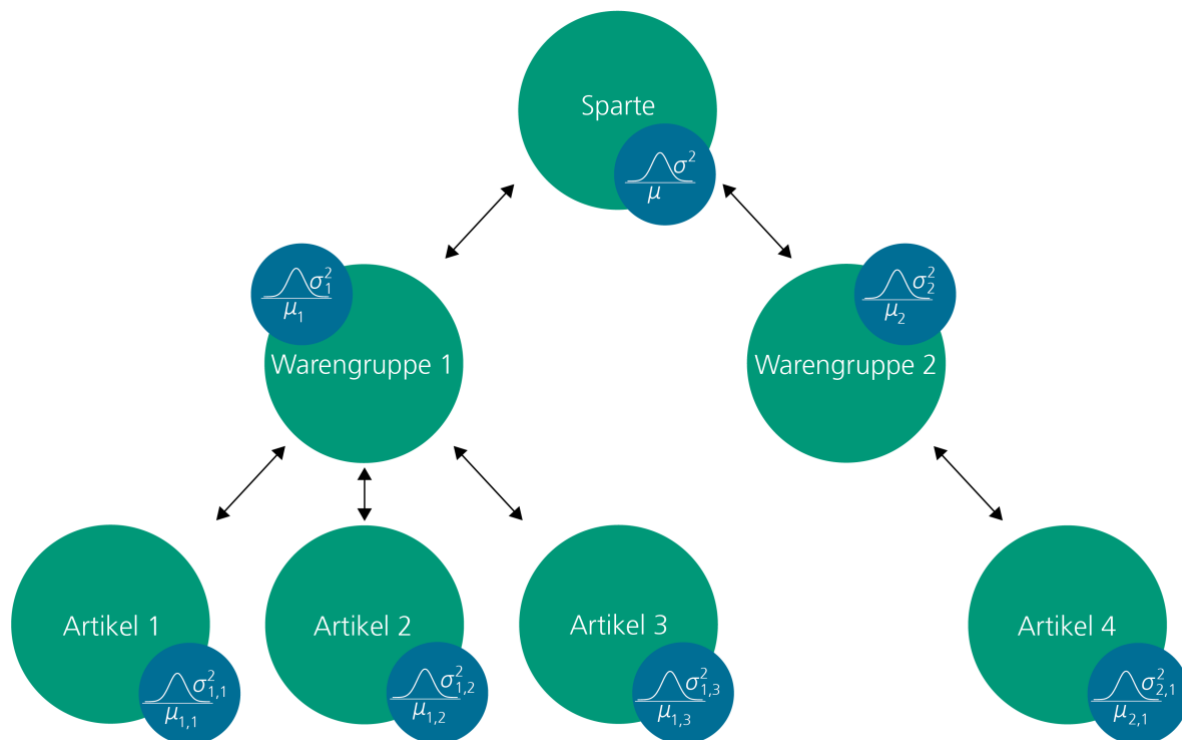


Abbildung 11: Darstellung eines hierarchischen Bayes-Modells. Die Verteilungsparameter der einzelnen Artikel hängen dabei von der jeweiligen Warengruppe ab, die Parameter der Warengruppen wiederum von der Sparte.

Zum Domänenwissen gehört oftmals ebenfalls ein Bewusstsein über Regeln, die allerdings meistens nur versprochen vorliegen (z. B. „Wenn wir eine Werbekampagne starten, steigt zwei Wochen später der Absatz des Produkts“). Die Idee von Neuro-Fuzzy ist, numerische Daten in Sprache zu übersetzen (Fuzzification), darauf dann die von Experten aufgestellten Regeln anzuwenden, und das Ergebnis schließlich wieder in Zahlen zu übersetzen (Defuzzification). Dies wird mit einer speziellen Architektur eines neuronalen Netzes erreicht (Stefan Siekmann et al. 1999). Die Wissenschaftler*innen der Arbeitsgruppe SCS erforschen hier, wie solche Architekturen flexibel in Code implementiert und problem-spezifisch adaptiert werden können.

4.2 QUANTIFIZIERUNG DER UNSICHERHEIT IN FORECASTING

Wie sorgfältig man ein Verfahren oder ein Modell auch auswählt und wie gut man es auch trainiert – eine perfekte Prognose wird nie möglich sein. Dies liegt zum einen daran, dass die Realität meistens komplizierter ist als alles, was modelliert werden kann. Es können nie alle potentiellen Einflussfaktoren und alle möglichen Formen der Abhängigkeit berücksichtigt werden. Bei statistischen Prognosemodellen nimmt man sogar im Allgemeinen an, dass man den Erwartungswert

vorhersagt, also „nur“ im Mittel richtigliegt. Hinzu kommt, dass auch dieser Erwartungswert aus den vorliegenden Daten geschätzt wird, wodurch immer ein Schätzfehler entsteht.

ML-Modelle hingegen können häufig so gut trainiert werden, dass sie die Trainingsdaten nahezu perfekt darstellen. Viele ML-Methoden, wie z. B. neuronale Netze, sind allerdings stark abhängig von der Startinitialisierung der Parameter. Je nachdem, mit welchen Parameterwerten der Algorithmus gestartet ist, können unterschiedliche Prognosen entstehen. Hier besteht die Unsicherheit der Prognose folglich darin, dass man nicht wissen kann, welches gelernte Modell das richtige ist, weil sie alle auf den bekannten Daten gleich gut sind.

Die Quantifizierung und Darstellung der Unsicherheit einer Prognose spielt für den Anwender eine große Rolle. Wird die Unsicherheit ganz ignoriert und lediglich eine Punktprognose ausgegeben, kann das zu Misstrauen in die Prognose führen, da sie ja immer wieder falsch liegt. Häufig ist außerdem eine Risikobewertung von Interesse. Beispielsweise möchte man beim Lagermanagement nicht nur den wahrscheinlichsten Bedarf vorhersagen, sondern auch einschätzen können, wie groß das Risiko ist, dass dieser prognostizierte Bedarf um eine bestimmte Menge über- oder unterschritten wird. So können unnötig große Sicherheitsbestände vermieden und gleichzeitig das Risiko, den Bedarf nicht decken zu können, minimiert werden. Auch bei der Preisprognose ist es von Vorteil, zusätzlich zur Punktprognose die Unsicherheit dieser Prognose miteinzubeziehen, um das Risiko eines Verlustes besser einschätzen zu können.

Um dieses Ziel zu erreichen, erforscht die Arbeitsgruppe SCS sogenannte Ensembleansätze, bei denen nicht nur die Prognose eines Modells genutzt wird, sondern die Prognosen vieler Modelle gruppiert werden (Ensemble). Liegen die Punktprognosen der Modelle nah beieinander, ist die Unsicherheit der Prognose klein und umgekehrt. Als Punktprognose des gesamten Ensembles wird eine Art Mittelwert aus den Prognosen der einzelnen Modelle gebildet (Lakshminarayanan et al. 2016). Wie robust diese Art der Prognoseunsicherheits-Bestimmung ist und inwiefern sie sich auf andere Ansätze als Ensemble von Neuronalen Netzen übertragen lässt, ist Gegenstand unserer Forschung.

Ein zweiter in der Forschung verbreitete Ansatz ist, die inhärenten Eigenschaften von Bayesianischen Methoden zu nutzen, um die Unsicherheit der Prognose zu quantifizieren. Da hier stets Verteilungen statt einzelner Parameter geschätzt werden, ist auch die Prognose immer eine Wahrscheinlichkeitsverteilung statt einer einfachen Punktprognose. Die Wissenschaftler*innen der Arbeitsgruppe SCS erforschen dabei, wie die beiden genannten Ansätze zur Bestimmung der Prognoseunsicherheit verglichen und kombiniert werden können.

Diese Verteilungsprognosen mit quantifizierter Unsicherheit können als Eingabe für Methoden der mathematischen Optimierungen unter Unsicherheiten, wie stochastische oder robuste Optimierung, genutzt werden. Anwendung hierfür ist zum Beispiel die Planung des Lagerbestandes, der auf Grund des unsicheren Absatzes, auf Basis von Prognosen optimiert werden muss. Die Arbeitsgruppe SCS erforscht gemeinsam mit der Friedrich-Alexander-Universität Erlangen Nürnberg das Zusammenspiel zwischen Machine Learning und Optimierung und entwickelt leistungsfähige Algorithmen für praxisrelevante Fragestellungen. Grundlage hierfür sind die mathematische Analyse der Problemstellungen und der Algorithmen.

4.3 KAUSALITÄTSANALYSEN

Die Entdeckung und das Beweisen von Kausalität sind schon seit langem ein Gegenstand der statistischen Forschung. Beispiele dafür sind die sogenannten Granger-Kausalität (Granger 1969) im Falle von Zeitreihenmodellen und die Arbeiten von Donald Rubin (Rubin 1974) und Judea Pearl (Pearl 2018). Allerdings entstehen auch hier neue Möglichkeiten durch größere Datenmengen und Rechenleistungen sowie der Verknüpfung mit ML-Methoden. Die Unterscheidung zwischen simpler Korrelation und tatsächlicher Kausalität ist alles andere als trivial. Als besonders aussagekräftig gelten Experimente wie randomisierte, kontrollierte Studien in der Medizin. Im Bereich der Supply Chain lassen sich solche Experimente im Allgemeinen allerdings nicht umsetzen. Unser Forschungsfokus liegt deshalb auf der Frage, wie man kausale Zusammenhänge durch die Analyse empirischer (Produktions-)Daten aufdecken kann. Für diesen Ansatz hat sich der Begriff „Causal Discovery“ durchgesetzt.

Dabei kann die Blickrichtung zum einen in die Vergangenheit gehen (Was war die Ursache dafür, dass Ereignis X eingetreten ist?), sodass man von Reverse Causal Inference spricht (Gelman 2011). Die zweite Möglichkeit ist die sogenannte Forward Causal Inference, bei der die Frage beantwortet werden soll: Was wird passieren, wenn X getan wird? Im Unterschied zu einer Prognose geht es hier allerdings nicht nur darum, den richtigen Wert vorherzusagen, sondern im Besonderen darum, eine kausale Erklärung für die beobachteten Abhängigkeiten zu finden. Erst durch diese kausale Erklärung kann robust prognostiziert werden, wie sich ein Sachverhalt unter einer Intervention verhalten wird; insbesondere dann, wenn die Voraussetzung von „unabhängigen und identisch verteilten Zufallsvariablen“ nicht fundiert angenommen werden kann. Es soll also das Verständnis über die Wirkzusammenhänge erweitert und so der Entscheidungsträger zu begründeten Entscheidungen befähigt werden. Außerdem sind die Erkenntnisse über kausale Zusammenhänge robuster als Prognosen und sogar auf andere Fragestellungen und Daten aus anderen Verteilungen übertragbar (Adversarial Vulnerability) (Schölkopf 2019).

Korrelationen können als gemeinsame Verteilungen interpretiert werden, die z. B. wie bei Bayes'schen Netzen gelernt werden können. Das Lernen von kausalen Zusammenhängen aus kausalen Netzen kann deshalb als Erweiterung dieses Ansatzes verstanden werden (Peters et al. 2017). Darstellmöglichkeiten dieser kausalen Netze sind beispielsweise Structural Causal Models (SCMs) und Causal Bayesian Networks.

Besteht bereits Domänenwissen oder Vermutungen über kausale Beziehungen, können diese in das Modell eingearbeitet werden. Schwieriger wird es, wenn die gesamte Struktur des Modells gelernt werden soll. Dazu gibt es statistische Methoden, die die Struktur des kausalen Graphen mithilfe von Tests auf bedingte Unabhängigkeit bestimmen, wie z. B. der IC-, der SGS- und der PC-Algorithmus (Peters et al. 2017). Ein weiterer Ansatz ist, denjenigen Graphen zu finden, dem für die vorliegenden Daten der höchste Wert z. B. der Likelihood oder der A-Posteriori-Dichte zugewiesen werden kann. Welche Graphenstrukturen bei dieser Suche berücksichtigt werden sollen, um möglichst schnell eine optimale Lösung zu finden, ist Gegenstand der Forschung.

Ein sehr aktueller Forschungsbereich ist die Anwendung von ML Methoden auf kausale Fragestellungen. Semi-überwachtes Lernen (also Lernen, bei dem für zusätzlichen

Informationsgewinn nicht nur annotiert, sondern auch nicht annotierte Daten verwendet werden) kann nach neueren Erkenntnissen für Reverse Causal Inference eingesetzt werden.

4.4 ONLINE-OPTIMIERUNG UND LERNENDE OPTIMIERUNGSVERFAHREN

Eine große, aktuelle Forschungsaufgabe liegt in der wechselseitigen Integration von KI-Lernverfahren und mathematischen Optimierungsverfahren. Einerseits liegen den Lernverfahren in aller Regel Optimierungsprobleme zu Grunde; man sucht die beste Ausgleichsfunktion, den besten Klassifikator, oder - ganz allgemein - den zur Erklärung der Daten besten funktionalen Zusammenhang zwischen Eingabe und Labels. Andersherum fand Entscheidungsunterstützung lange Zeit als rein sequentieller Prozess zwischen Statistik/KI und Optimierung statt. Zur Schätzung der Parameter in einem Optimierungsproblem wurden zum Beispiel Regressionsverfahren genutzt, um zukünftige Trends vorherzusagen. Nach Schätzung dieser Parameter fand auf deren Grundlage die Optimierung statt. Zukünftig wird es von entscheidender Bedeutung sein, Algorithmen für Optimierungsprobleme zu entwickeln, die auf sich dynamisch verändernden Daten basieren, wie sie zum Beispiel in der Echtzeitoptimierung vorkommen. Hier genügt es nicht mehr, die Laufzeit der eingesetzten Algorithmen klein zu halten (was für sich genommen schon eine Herausforderung ist). Es müssen auch Gütegarantien für die berechneten Lösungen unter zeitlich veränderlichen Daten eingehalten werden.

Für die oftmals hochdynamischen Randbedingungen von Logistikketten wird die Verfügbarkeit solcher Verfahren von besonderer Wichtigkeit sein. In einem gemeinsamen Forschungsprojekt zwischen der Arbeitsgruppe SCS und der Friedrich-Alexander-Universität Erlangen-Nürnberg wurde beispielsweise gezeigt, dass es möglich ist, Prioritäten und Kostenfunktionen in Planungsproblemen aus beobachteten Entscheidungen der Vergangenheit abzuleiten. Dieser methodische Ansatz für logistische Fragestellungen ist insbesondere deshalb von hoher Relevanz, da viele Planungsregeln implizit, im Sinne von nicht formalisiert und der Automatisierung schwer zugänglich, sind. Die Fähigkeit, aus vergangenen Entscheidungen explizite Planungsregeln herzuleiten, sorgt so für einen höheren Automatisierungsgrad der Logistikkette und die Objektivierbarkeit logistischer Planungsentscheidungen. Die Weiterentwicklung dieses Ansatzes im Rahmen lernender Optimierungsverfahren wird in der zukünftigen Forschung einen besonderen Stellenwert einnehmen. Der praktische Einsatz dieser Verfahren im industriellen Kontext wird für den Anwender einen entscheidenden Wettbewerbsvorteil ausmachen.

5 ANHANG

5.1 LITERATUR

5.2 ABBILDUNGSVERZEICHNIS

Abbildung 1: Aufbau und Angebot des ADA Lovelace Centers - Vernetzung von Industrie und Forschung	8
Abbildung 1: Team Data Science Process von Microsoft	13
Abbildung 2: Graphische Darstellung eines Entscheidungsbaumes	13
Abbildung 3: Graphische Darstellung Support Vector Machine	14
Abbildung 4: Graphische Darstellung eines neuronalen Netzes mit zwei versteckten Schichten	15
Abbildung 5: Graphische Darstellung eines Bayes'schen Netzes	16
Abbildung 6: Graphische Darstellung eines Error-Correctional Neural Networks, bei dem der Fehler aus dem vorangegangenen Zeitschritt für die Vorhersage des nächsten Zeitschrittes verwendet wird.	18
Abbildung 7: Graphische Darstellung einer Clusteranalyse	19
Abbildung 8: Graphische Darstellung einer Hauptachsentransformation. Die Hauptkomponenten sind orthogonal zueinander und damit unkorreliert.	21
Abbildung 9: Graphische Darstellung eines Branch-and-Bound Baumverfahren	24
Abbildung 10: Graphische Darstellung eines Konvergenz-Optimierungsverfahren	25

5.3 TABELLENVERZEICHNIS

Tabelle 1: Tabelle (Quelle: Eigene Darstellung)	3
---	---

5.4 DIE AUTOREN



Dr. Andreas Bärman ist Wissenschaftler am Lehrstuhl für Analytics & Mixed-Integer Optimization der Friedrich-Alexander-Universität Erlangen-Nürnberg. In zahlreichen Forschungs- und Industrieprojekten forscht er gemeinsam mit der Arbeitsgruppe SCS an der Theorie und praktischen Lösung von herausfordernden diskreten Optimierungsproblemen.

Im ADA Lovelace Center leitet er die Forschungsarbeiten im Bereich mathematische Optimierung. Die in seinem Team entwickelten Techniken fließen ein in Softwaretools zur Entscheidungsunterstützung in den Branchen Logistik, Produktion, Mobilität und Energie.



Julius Mehringer ist wissenschaftlicher Mitarbeiter bei der Fraunhofer-Arbeitsgruppe SCS und forscht in der Gruppe Data Science zum Einsatz von künstlicher Intelligenz in Produktionsumgebungen. Seine Schwerpunkte liegen dabei auf bayesianischen Modellen zur Vorhersage, deren Verknüpfung mit der mathematischen Optimierung sowie auf Analysen zur Identifikation von kausalen Effekten. Er hat u.a. International Information Systems Management und Survey Statistik studiert. Im ADA Lovelace Center leitet er das Projekt ALONE.



Christian Menden leitet seit 2019 die Abteilung Analytics der Fraunhofer-Arbeitsgruppe SCS mit ihren vier Gruppen „Process Intelligence“, „Data Science“, „Optimization“ sowie „Data Efficient and Automated Learning“ mit rund 30 Forscher*innen. Seine Promotion behandelte die Anwendung von Data Augmentation Methoden, um maschinelle Lernverfahren im Supply Chain Management bei schwieriger bzw. unzureichender Datenlage einzusetzen. Seit 2017 beschäftigt sich der Wissenschaftler mit Data-Analytics Fragestellungen im Supply Chain Management und ist als Lehrbeauftragter an der TH Nürnberg, der FH Würzburg-Schweinfurt sowie der Otto-Friedrich-Universität Bamberg tätig.



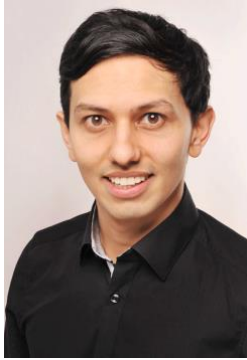
Dr. Ursula Neumann ist Gruppenleiterin der Gruppe Data Science in der Fraunhofer-Arbeitsgruppe für Supply Chain Services des Fraunhofer IIS.

Nach ihrem Studium der Mathematik an der Universität Regensburg promovierte Ursula Neumann im Bereich Biostatistics zum Thema Ensemble Feature Selection für binäre Klassifikationen an der Technischen Universität München.

Das Forschungsinteresse von Ursula Neumann betrifft die Entwicklung von KI-basierten Methoden für die Anwendung in Handel, Logistik und Produktion.



Julia Schemm ist wissenschaftliche Mitarbeiterin bei der Fraunhofer-Arbeitsgruppe SCS und forscht in der Gruppe Data Science zum Einsatz von künstlicher Intelligenz im Großhandel. Ihr besonderer Fokus liegt dabei auf Bedarfsprognosen mithilfe Neuronaler Netze und bayesianischer Methoden. Sie hat Mathematik in Bonn und Statistische Wissenschaften in Bielefeld studiert.



Oskar Schneider ist wissenschaftlicher Mitarbeiter der Fraunhofer-Arbeitsgruppe SCS in der Gruppe Optimization und Doktorand an der Friedrich-Alexander-Universität Erlangen-Nürnberg am Lehrstuhl für Analytics & Mixed-Integer Optimization.

Er forscht im ADA Lovelace Center zum Einsatz von künstlicher Intelligenz in Produktionsumgebungen. Sein Schwerpunkt liegt dabei auf Optimierungsmodellen für die Planung der Produktion von Lebensmitteln.

Er hat Technomathematik an der Friedrich-Alexander-Universität Erlangen-Nürnberg studiert.



Benedikt Sonnleitner ist wissenschaftlicher Mitarbeiter bei der Fraunhofer-Arbeitsgruppe SCS und forscht in der Gruppe Data Science zum Einsatz von Künstlicher Intelligenz in der Logistik. Sein Schwerpunkt liegt dabei auf Zeitreihenprognose, insbesondere für Bedarfe. Daneben leitet er verschiedene Forschungsprojekte zum Einsatz von Machine-Learning-Methoden in der Logistik. Er hat Wirtschaftsingenieurwesen und Logistik mit dem Schwerpunkt auf Operations Research studiert.



Markus Weissenböck ist Leiter der Gruppe Optimization bei der Fraunhofer-Arbeitsgruppe für Supply Chain Services des Fraunhofer IIS, die unter anderem Methoden und Verfahren für anwendungsnahe Data Analytics und mathematische Optimierungsverfahren im Produktions- und Logistikumfeld entwickelt. Er ist Experte für datenbasierte Entscheidungsunterstützungssysteme.

6 LITERATURVERZEICHNIS

- Aggarwal, Charu C. (2018): Neural Networks and Deep Learning. A Textbook. Cham: Springer International Publishing.
- Applegate, David L.; Bixby, Robert E.; Chvátal, Vasek; Cook, William J. (2006): The Traveling Salesman Problem. A Computational Study: Princeton University Press.
- Bouveyron, Charles; Côme, Etienne; Jacques, Julien (2015): The discriminative functional mixture model for a comparative analysis of bike sharing systems. In: *Ann. Appl. Stat.* 9 (4), S. 1726–1760. DOI: 10.1214/15-AOAS861.
- Box, G. E. P.; Jenkins, G. M. (1970): Time series analysis: Forecasting and control: San Francisco: Holden-Day.
- Breiman, Leo (2001): Random Forests. In: *Machine Learning* 45 (1), S. 5–32. DOI: 10.1023/A:1010933404324.
- Brown, Robert Goodell (1959): Statistical forecasting for inventory control: McGraw/Hill.
- Cristianini, Nello; Shawe-Taylor, John (2012): An introduction to support vector machines and other kernel-based learning methods. 13. print. Cambridge: Cambridge Univ. Press.
- Davis, L. (1991): Handbook of genetic algorithms. Online verfügbar unter <http://papers.cumincad.org/cgi-bin/works/paper/eaca>.
- Dorigo, Marco; Di Caro, Gianni; Gambardella, Luca M. (1999): Ant Algorithms for Discrete Optimization. In: *Artificial Life* 5 (2), S. 137–172. DOI: 10.1162/106454699568728.
- Elias Gyftodimos; Peter A. Flach (2002): Hierarchical Bayesian Networks: A Probabilistic Reasoning Model for Structured Domains. Online verfügbar unter https://www.researchgate.net/publication/2530728_Hierarchical_Bayesian_Networks_A_Probabilistic_Reasoning_Model_for_Structured_Domains.
- Ester, Martin; Kriegel, Hans-Peter; Sander, Jiirg; Xu, Xiaowei (1996): A Density-Based Algorithm for Discovering Clusters in Large Spatial Databases with Noise. Proceedings ; [August 2 -4, 1996, Portland, Oregon]. In: *KDD-96 Proceedings* (34).
- Gelman, Andrew (2011): Causality and Statistical Learning. In: *American Journal of Sociology* 117 (3), S. 955–966. DOI: 10.1086/662659.
- Gordon, A. D.; Breiman, L.; Friedman, J. H.; Olshen, R. A.; Stone, C. J. (1984): Classification and Regression Trees. In: *Biometrics* 40 (3), S. 874. DOI: 10.2307/2530946.
- Guyon, Isabelle; Elisseeff, André (2003): An Introduction to Variable and Feature Selection. In: *Journal of Machine Learning Research* 3 (Mar), S. 1157–1182.
- Hochreiter, S.; Schmidhuber, J. (1997): Long short-term memory. In: *Neural computation* 9 (8), S. 1735–1780. DOI: 10.1162/neco.1997.9.8.1735.
- Holt, Charles C. (2004): Forecasting seasonals and trends by exponentially weighted moving averages. In: *International Journal of Forecasting* 20 (1), S. 5–10. DOI: 10.1016/j.ijforecast.2003.09.015.

- James, Gareth; Witten, Daniela; Hastie, Trevor; Tibshirani, Robert (2017): An introduction to statistical learning. With applications in R. Corrected at 8th printing. New York, Heidelberg, Dordrecht, London: Springer (Springer texts in statistics, 103).
- Johnson, S. C. (1967): Hierarchical clustering schemes. In: *Psychometrika* 32 (3), S. 241–254. DOI: 10.1007/BF02289588.
- Lakshminarayanan, Balaji; Pritzel, Alexander; Blundell, Charles (2016): Simple and Scalable Predictive Uncertainty Estimation using Deep Ensembles. Online verfügbar unter <https://arxiv.org/pdf/1612.01474>.
- Lawler, E. L.; Wood, D. E. (1966): Branch-and-Bound Methods: A Survey. In: *Operations Research* 14 (4), S. 699–719. DOI: 10.1287/opre.14.4.699.
- Lloyd, S. (1982): Least squares quantization in PCM. In: *IEEE Trans. Inform. Theory* 28 (2), S. 129–137. DOI: 10.1109/TIT.1982.1056489.
- Nelder, J. A.; Wedderburn, R. W. M. (1972): Generalized Linear Models. In: *Journal of the Royal Statistical Society. Series A (General)* 135 (3), S. 370–384. DOI: 10.2307/2344614.
- Padberg, M.; Rinaldi, G. (1987): Optimization of a 532-city symmetric traveling salesman problem by branch and cut. In: *Operations Research Letters* 6 (1), S. 1–7. DOI: 10.1016/0167-6377(87)90002-2.
- Pearl, J. (1985): Bayesian Networks: A Model of Self-activated Memory for Evidential Reasoning: UCLA, Computer Science Department. Online verfügbar unter <https://books.google.de/books?id=1sfMOgAACAAJ>.
- Peter J. M. van Laarhoven; Emile H. L. Aarts (1987): Simulated annealing. In: *Simulated Annealing: Theory and Applications*: Springer, Dordrecht, S. 7–15. Online verfügbar unter https://link.springer.com/chapter/10.1007/978-94-015-7744-1_2.
- Peters, Jonas; Janzing, Dominik; Schölkopf, Bernhard (2017): Elements of causal inference. Foundations and learning algorithms. Cambridge, Massachusetts, London, England: MIT Press (Adaptive computation and machine learning).
- Rego, César; Gamboa, Dorabela; Glover, Fred; Osterman, Colin (2011): Traveling salesman problem heuristics: Leading methods, implementations and latest advances. In: *European Journal of Operational Research* 211 (3), S. 427–441. DOI: 10.1016/j.ejor.2010.09.010.
- Schölkopf, Bernhard (2019): Causality for Machine Learning. Online verfügbar unter <https://arxiv.org/pdf/1911.10500>.
- Schulte, Christof (2009): Logistik. Wege zur Optimierung der Supply Chain. 5., überarb. und erw. Aufl. München: F. Vahlen (Vahlens Handbücher der Wirtschafts- und Sozialwissenschaften).
- Stefan Siekmann; R. Neuneier; H. G. Zimmermann; R. Kruse (1999): Neuro Fuzzy Systems for Data Analysis. In: *undefined*. Online verfügbar unter <https://www.semanticscholar.org/paper/Neuro-Fuzzy-Systems-for-Data-Analysis-Siekmann-Neuneier/e9d01cd5d5d078488cb17e3d374298efa7380a76>.
- Wang, Gang; Gunasekaran, Angappa; Ngai, Eric W.T.; Papadopoulos, Thanos (2016): Big data analytics in logistics and supply chain management: Certain investigations for research and applications. In: *International Journal of Production Economics* 176, S. 98–110. DOI: 10.1016/j.ijpe.2016.03.014.

What is the Team Data Science Process? (2016). Online verfügbar unter <https://docs.microsoft.com/en-us/azure/machine-learning/team-data-science-process/overview>, zuletzt geprüft am 26.02.2021.

Winters, Peter R. (1960): Forecasting Sales by Exponentially Weighted Moving Averages. In: *Management Science* 6 (3), S. 324–342. DOI: 10.1287/mnsc.6.3.324.

Wold, Svante; Esbensen, Kim; Geladi, Paul (1987): Principal component analysis. In: *Chemometrics and Intelligent Laboratory Systems* 2 (1-3), S. 37–52. DOI: 10.1016/0169-7439(87)80084-9.

Zimmermann, Hans-Georg; Grothmann, Ralph; Tietz, Christoph (2012): Forecasting Market Prices with Causal-Retro-Causal Neural Networks. In: Diethard Klatte, Hans-Jakob Lüthi und Karl Schmedders (Hg.): Operations research proceedings 2011. Selected papers of the International Conference on Operations Research (OR 2011), August 30 - September 2, 2011, Zurich, Switzerland. Berlin, Heidelberg, 2012. Berlin, Heidelberg: Springer Berlin Heidelberg (Operations Research Proceedings, GOR (Gesellschaft für Operations Research e.V.)), S. 579–584.

Zimmermann, Hans-Georg; Grothmann, Ralph; Tietz, Christoph; Jouanne-Diedrich, Holger von (2011): Market Modeling, Forecasting and Risk Analysis with Historical Consistent Neural Networks. In: Bo Hu, Karl Morasch, Stefan Pickl und Markus Siegle (Hg.): Operations Research Proceedings 2010. Selected Papers of the Annual International Conference of the German Operations Research Society. Berlin, Heidelberg, 2011. Berlin, Heidelberg: Springer-Verlag Berlin Heidelberg (Operations Research Proceedings, GOR (Gesellschaft für Operations Research e.V.)), S. 531–536.

Zimmermann, Hans-Georg; Neuneier, Ralph; Grothmann, Ralph (2002): Modeling Dynamical Systems by Error Correction Neural Networks. In: Abdol S. Soofi und Liangyue Cao (Hg.): Modelling and Forecasting Financial Data. Techniques of Nonlinear Dynamics. New York: Springer Science+Business Media LLC (Studies in Computational Finance, 2), S. 237–263. Online verfügbar unter https://doi.org/10.1007/978-1-4615-0931-8_12.

Fraunhofer-Institut für Integrierte Schaltungen IIS

Institutsleitung

Prof. Dr.-Ing. Albert Heuberger (geschäftsführend)

Prof. Dr.-Ing. Bernhard Grill

Prof. Dr. Alexander Martin

Am Wolfsmantel 33

91058 Erlangen

Telefon +49 9131 776-0

info@iis.fraunhofer.de

www.iis.fraunhofer.de

Fraunhofer-Arbeitsgruppe für Supply Chain Services des Fraunhofer IIS

Leitung

Prof. Dr. Alexander Pflaum

Geschäftsführung

Dr.-Ing. Roland Fischer

Nordostpark 93

90411 Nürnberg

Kontakt

Telefon +49 911 58061-9500

info@scs.fraunhofer.de

www.scs.fraunhofer.de