

INHALT

01	Management Summary	2
02	Willkommen im KI-Kosmos	3
	2.1 Einordnung der wichtigsten Begriffe und Disziplinen	3
	2.1.1 Künstliche Intelligenz	3
	2.1.2 Data Science	3
	2.1.3 Machine Learning	4
	2.1.4 Deep Learning & künstliche neuronale Netze	4
	2.2 Wie funktioniert Machine Learning?	5
03	Auf Machine-Learning-Erkundungsmission:	7
	3.1 Identifizieren sinnvoller Use Cases im Unternehmen	7
	3.2 Auswahl eines geeigneten Machine-Learning-Ansatzes	8
	3.2.1 Supervised Learning: Daten sind gelabelt	8
	3.2.2 Unsupervised Learning: Daten sind nicht gelabelt	9
	3.2.3 Semi-Supervised Learning: Daten sind teilweise gelabelt	10
	Exkurs I: Erfolgreiche Accso KI-Mission:	
	Recommendation-System für ein Onlineportal	12
	Exkurs II: Erfolgreiche Accso KI-Mission:	
	Proof of Concept (PoC) zur Voraussfüllung von Buchungen in einer Buchhaltungssoftware	14
04	Takeoff, Kurskorrektur und Satellitennetz: Von der ersten ML-Lösung bis zum Betrieb der Anwendung	16
	4.1 Entwicklung einer Machine-Learning-Lösung	16
	Exkurs III: Was kostet die Entwicklung?	21
	4.2 Betrieb und Weiterentwicklung einer Machine-Learning-Lösung	23
05	Aus dem Accsonauten-Alltag: Erfolgsfaktoren und Lessons Learned aus der Praxis	25
	5.1 Erfolgsfaktoren für ein KI-Projekt	25
	5.1.1 Technische Aspekte	25
	5.1.2 Organisatorische Faktoren	26
	5.1.3 Einsatz in der Produktion	27
	5.2 Dos and Don'ts bei der Umsetzung von KI-Projekten	28
	5.2.1 Dos	28
	5.2.2 Don'ts	32
06	Weiterführende Literatur	33

01

Management Summary

„Wenn ich es hinbekomme, dass wir unser Programm durchführen und dass wir als Freunde zurückkommen, dann ist das für mich eine großartige Mission geworden.“

Der deutsche Astronaut Alexander Gerst vor seiner zweiten Weltraummission.

In unser aller Leben spielen Künstliche Intelligenz (KI) und Machine Learning (ML) eine Rolle – oft auch ohne, dass wir uns dessen bewusst sind. Viele Innovationen der letzten Jahre lassen sich auf Machine Learning und Künstliche Intelligenz zurückführen. So werden für die Routenberechnung in Navigationsapps, für Sprachassistenten im Smartphone oder in den Filmempfehlungen unseres Streamingdienstes Machine Learning Techniken angewandt.

Das Potenzial erstreckt sich dabei nicht nur auf große Technikfirmen und Konzerne. Unternehmen jeglicher Größe können die Möglichkeiten von Machine Learning und Künstlicher Intelligenz nutzen. In dem vorliegenden Whitepaper geht es genau darum: Es nimmt Unternehmen mit auf KI-Erkundungsmission, damit sie sich mit den zugrundeliegenden Methoden und Projektvorgehen vertraut machen können. Ob aus einer Mission eventuell eine andauernde Freundschaft wird, kann jedes Unternehmen selbst entscheiden.

Mit allerhand praktischer Erfahrung aus dem Accsonauten-Alltag loten wir in diesem Whitepaper die Einsatzmöglichkeiten von Machine Learning für Unternehmen mit folgenden Fragestellungen aus:

- Was sind die wichtigsten Grundbegriffe rund um Machine Learning?
- Wie funktioniert die Technologie?
- Welche Bereiche eignen sich für den Einsatz von Machine Learning?
- Wie läuft ein Machine Learning Projekt ab?
- Was kostet ein Machine Learning Projekt und der Betrieb der Anwendung?
- Was sind entscheidende Erfolgsfaktoren?

02

Willkommen im KI-Kosmos

2.1 Einordnung der wichtigsten Begriffe und Disziplinen

Künstliche Intelligenz, Data Science und Machine Learning werden oft synonym verwendet. Die Fachausdrücke sind eng miteinander verbunden, beschreiben jedoch unterschiedliche Aspekte. Ganz nach dem sokratischen Spruch „Der Beginn der Weisheit liegt in der Definition der Begriffe“ widmen wir uns zu Beginn unserer Machine-Learning-Mission den wichtigsten Termini.

2.1.1 Künstliche Intelligenz

Definition von Künstlicher Intelligenz nach dem Netzwerk Mittelstand digital:

„Angelehnt an die Leistungsfähigkeit der menschlichen Intelligenz fokussiert Künstliche Intelligenz die Lösung konkreter Anwendungsprobleme und unterstützt den Menschen bei Arbeits- und Entscheidungsprozessen. Kennzeichnend für ein KI-System ist die Lernfähigkeit auf Basis von Daten.“

Künstliche Intelligenz vereint als Überbegriff **sämtliche Ansätze, die Maschinen entwickeln, die eine Form von Intelligenz zeigen bzw. sich intelligent verhalten.** Im Bereich der geschäftlich genutzten KI ist dabei eigentlich immer die sogenannte schwache KI gemeint. Hierbei handelt es sich um eine spezialisierte KI, die nur im Kontext ihrer Anwendung intelligentes Verhalten zeigt. Hingegen verhält sich eine starke KI in jeder Situation intelligent – unabhängig von der Problematik oder dem Themengebiet. Obwohl in den letzten zwei Jahrzehnten viele Fortschritte in der KI-Entwicklung erfolgten, sind wir von einer starken KI noch sehr weit entfernt, auch wenn Sci-Fi-Filme heutzutage oft ein anderes Bild vermitteln.

2.1.2 Data Science

Data Science definiert ein Teilgebiet Künstlicher Intelligenz. Es analysiert und verarbeitet Daten, um **datengetriebene KI-Modelle** zu **entwickeln.** Im Gegensatz zu regelbasierten KI-Modellen können datengetriebene KI-Modelle eigenständig Zusammenhänge und Regeln aus vorhandenen Daten ableiten. Dies hat zum einen den Vorteil,

dass über den modellierten Zusammenhang wenig bis kein Vorwissen vorhanden sein muss, zum anderen aber auch, dass die KI stetig aus zusätzlichen Daten weiter lernen und sich verbessern kann.

Damit unterscheidet sich Data Science auch von Business Intelligence (BI), welche sich zwar ebenfalls mit der Analyse von Daten beschäftigt, jedoch Daten hinsichtlich verschiedener statistischer Kenngrößen auswertet. BI umfasst auch klassische Data-Engineering-Aufgaben, wie den Aufbau und den Betrieb einer geeigneten Datenplattform. Für die Entwicklung und den langfristigen Betrieb einer datengetriebenen KI-Anwendung sind diese genauso wichtig wie die KI selbst, werden jedoch im Kontext dieses Whitepapers als vorhanden vorausgesetzt.

2.1.3 Machine Learning

Machine Learning beschäftigt sich mit der **Entwicklung verschiedener Algorithmen und Werkzeuge, mit denen Computer eigenständig Zusammenhänge, Muster und Regeln aus Daten lernen können**. Auf Basis des Gelernten leiten Algorithmen dann Entscheidungen für neue Daten und Vorhersagen ab. Machine Learning liefert damit die Softwaregrundlage, die im Kontext von Data Science zur Entwicklung von KI-Modellen eingesetzt wird.

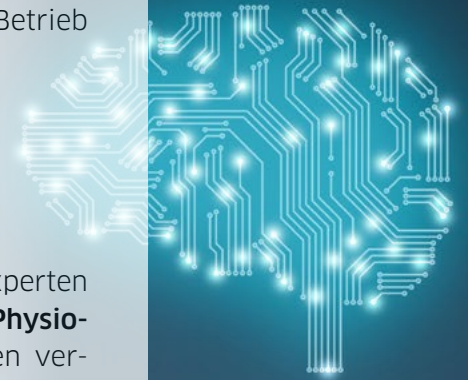
Neben der reinen Entwicklung geht es uns beim Machine Learning aber auch darum, zu wissen, wie diese Lösungen in die Produktion gebracht werden können. Dazu gehören neben dem Aufbau

einer Infrastruktur auch das Ausrollen, der Betrieb und die Überwachung von KI-Modellen.

2.1.4 Deep Learning & künstliche neuronale Netze

Bereits in den 1930er Jahren wollten Experten und Expertinnen herausfinden, wie die **Physiologie des menschlichen Gehirns** mit seinen verbundenen Neuronen **mathematisch beschrieben** und **künstlich nachgebildet** werden kann. Durch die begrenzten Hardwareressourcen der Zeit dauerte es aber bis in die 2000er Jahre, bis tiefe Netzarchitekturen bestehend aus vielen Schichten als künstliche neuronale Netze (KNN) effektiv einsetzbar waren. Im Unterschied zu den flachen Architekturen der frühen Jahre wird der Einsatz moderner neuronaler Netze daher auch unter dem Begriff Deep Learning zusammengefasst.

KI und Machine Learning gewannen gerade in den letzten Jahren durch entscheidende Weiterentwicklungen im Bereich der Bild- und Sprachverarbeitung sowie durch die Verfügbarkeit von geeigneter Hardware eine enorme Popularität. Dieser Erfolg ist zum Großteil auf Deep Learning zurückzuführen.



2.2 Wie funktioniert Machine Learning?

Im Gegensatz zu Modellen klassischer KI-Verfahren bildet ein Machine-Learning-Modell keine Sammlung von Entscheidungsregeln, sondern basiert auf einem selbstlernenden Algorithmus. Wird der Algorithmus in der Trainingsphase mit Daten gefüttert, so kann er diese selbstständig analysieren, Zusammenhänge erkennen und daraus Regeln ableiten, um das vorgegebene Ziel zu erreichen.

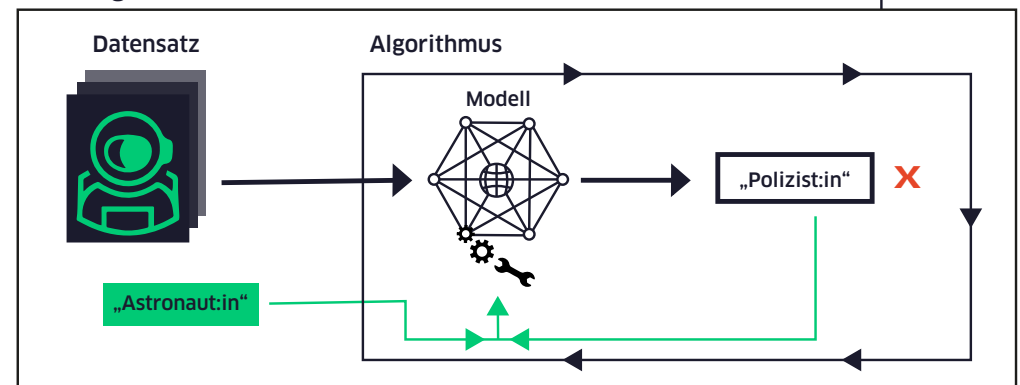
Ein häufiges Ziel besteht z. B. darin, möglichst wenige Fehler oder eine möglichst hohe Genauigkeit bei einer Vorhersage zu erreichen. Wie genau dieses Ziel gestaltet ist, hängt aber auch vom gewählten Machine-Learning-Ansatz ab. **Die so gelernten Zusammenhänge und Regeln bilden zusammen mit dem Algorithmus das Modell.** Dieses kann dann eingesetzt werden, um das Gelernte auf neue Daten anzuwenden. Dieser Schritt wird auch als Inferenz bezeichnet.

Generell können die Daten in unterschiedlicher Form vorliegen: Tabellen, CSV-Dateien, Bilder, Texte, Audio usw. Oft sind die einzelnen Datenpunkte jeweils mit einem Label verknüpft. Das ist der Wert der Zielgröße, die das Machine-Learning-Modell vorhersagen soll.

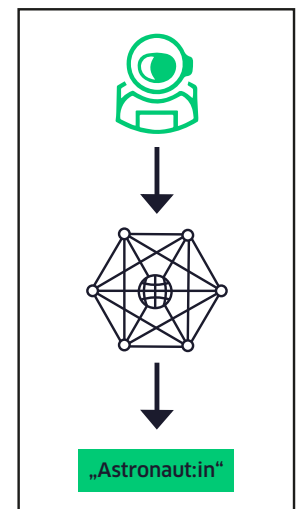
Nehmen wir ein Beispiel: Unser Modell soll aus einer Vielzahl von Bildern, die Menschen und ihre Berufe zeigen, Fotos mit Astronautinnen und Astronauten herausfiltern. In den Trainings-

daten werden die Astronautenfotos mit dem Label „Astronaut:in“ versehen. Das System kennt jetzt die Zielgröße und untersucht alle weiteren Fotos auf Gemeinsamkeiten hinsichtlich der Zielbilder. So achtet es bei den Bildern z.B. auf die menschliche Statur, einen runden Helm oder auch die Farbe der Kleidung. Welche Unterscheidungsmerkmale sind relevant? Das ist die entscheidende Frage. Die meisten Daten sind typischerweise irrelevant oder redundant und helfen dem Algorithmus nicht, sein Ziel zu erreichen, möglichst wenige Fehler bei der Unterscheidung zwischen Fotos mit und ohne Astronautin oder Astronaut zu machen. Ein Helm im Bild bedeutet nicht automatisch, dass es sich um eine Astronautin oder einen Astronauten handelt. Es könnte auch eine Rennfahrerin oder Feuerwehrmann dargestellt sein. Auch Eigenschaften wie Augen- und Haarfarbe sind irrelevant. Die möglichen Unterscheidungsmerkmale bezeichnet man als Features.

Training



Inferenz



Trainiertes Modell

Funktionsweise von Machine Learning

Neben dem Label nutzt der Algorithmus diese, um die richtigen Zielwerte zu erreichen. Die Auswahl der relevanten Features kann entweder auf Basis von Fachwissen manuell oder von manchen Algorithmen auch selbst durchgeführt werden.

Während des Trainingsprozesses erhält der Machine-Learning-Algorithmus iterativ Trainingsdaten zugeführt, wodurch das Modell Form annimmt. Auf Basis des bereits Gelernten kann das Modell den Input verarbeiten und liefert eine Prognose. Abhängig von der Qualität der Prognose passt der Algorithmus die gelernten Regeln im Modell an. Im Trainingsprozess kann es leicht dazu kommen, dass die präsentierten Daten einfach auswendig gelernt werden, um eine möglichst hohe Trainingsgenauigkeit zu erreichen.

Um dies zu vermeiden, setzen die Fachleute gesonderte Validierungs- bzw. Testdatensätze ein, welche dem Lernprozess vorenthalten werden.

Lernfähige Machine-Learning-Algorithmen sind statistische Verfahren. Das heißt, sie **lernen aus den Daten die wahrscheinlichsten Regeln** und **liefern** nach abgeschlossenem Training **die wahrscheinlichste Entscheidung**. Damit eine Machine-Learning-Anwendung Ergebnisse mit sehr hoher Genauigkeit liefern kann, sind der Umfang und die Qualität der Trainingsdaten mindestens genauso wichtig wie die Auswahl des passenden Algorithmus und das daraus resultierende Modell.

Bevor nun über geeignete Lösungsansätze nachgedacht werden kann, gilt es die viel zentralere Frage zu klären: Ist Machine Learning überhaupt das geeignete Mittel, um die selbst gesteckten Ziele im Kontext des geplanten Use Case zu erreichen?

3.1 Identifizieren sinnvoller Use Cases im Unternehmen

Form follows function! Dieser bekannte Designleitsatz gilt auch beim Einsatz von Machine Learning in Unternehmen. Im Zentrum aller Überlegungen liegt die zu erfüllende Aufgabe. Das gesuchte Design muss dieser Aufgabe gerecht werden und nicht andersherum.

Allzu oft erleben wir, dass die Methode oder Technologie im Vordergrund steht und dann passende Herausforderungen geschaffen werden. Dies trifft vor allem in unserer Branche auf Trendthemen wie Machine Learning oder Blockchain zu. Oftmals gehen damit dann auch unrealistische und überzogene Erwartungen einher. Grundsätzlich gilt: Machine Learning kann nicht alle Probleme lösen.

Wir gehen davon aus, dass heutige KI (und auch die in absehbarer Zukunft) nur für spezifische Anwendungsfälle gut geeignet ist. Diese wiederum lassen sich dann mit teilweise erstaunlicher Leistung bewältigen. Was zeichnet diese Anwendungsfälle, in denen der Einsatz von Machine Learning

angebracht und gewinnbringend sein kann, aktuell aus? Die folgenden Aspekte können einen ersten Ansatz darstellen:

- Zu hohe Komplexität für eine manuelle Programmierung der Lösung (z. B. individuelle Empfehlungen, Bilderkennung)
- Eine regelbasierte Lösung ist denkbar, jedoch sind die Logik und das resultierende Regelwerk zu unübersichtlich und komplex
- Hoher monotoner Aufwand
- Nennenswerte Fehleranfälligkeit bei der manuellen Ausführung
- Vorhersagen, die auf klassischen statistischen Verfahren basieren
- Bestehende Zusammenhänge in Daten entdecken und nutzen (z. B. „Was macht die größten Kundengruppen aus?“)

Wenn Unternehmen ähnliche Aspekte in ihren Aufgaben identifiziert haben, lohnt es sich, über den Einsatz von Machine Learning in diesen Bereichen nachzudenken.

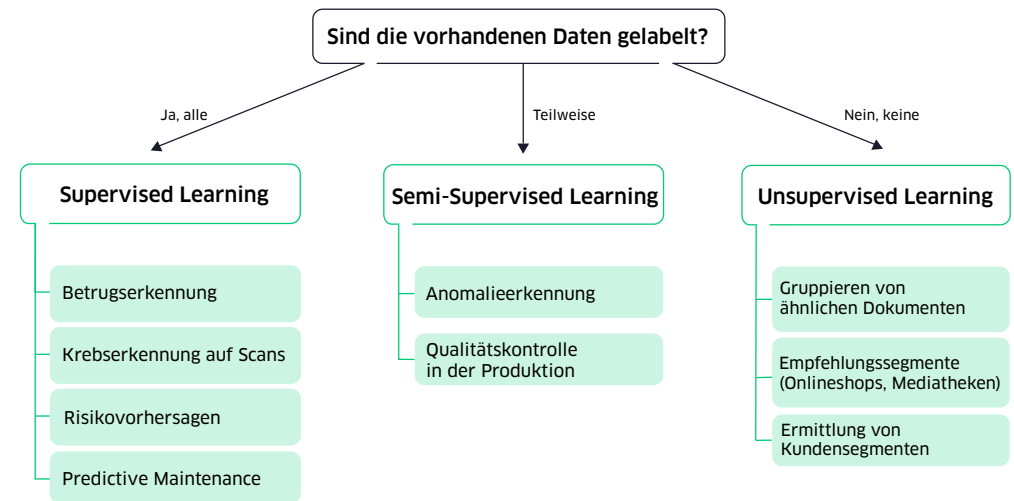
3.2 Auswahl eines geeigneten Machine-Learning-Ansatzes

Form follows function – again! Im nächsten Schritt wollen wir uns mit der Wahl des geeigneten Machine-Learning-Verfahrens beschäftigen. Ein tieferes Verständnis des Anwendungsfalls und die richtige Wahl des Lernverfahrens haben eine enorme Auswirkung auf die Erfolgschancen. Die **Festlegung des Trainingsziels** bereitet **die größten Schwierigkeiten**, nicht unbedingt die Auswahl des Algorithmus selbst. Das haben wir in unseren Projekten festgestellt. Daher muss die zentrale Fragestellung lauten: „Sind die vorhandenen Daten gelabelt, also sind die Werte der Zielgröße für jeden Datenpunkt bekannt?“

Gemäß den Antwortmöglichkeiten auf diese Frage gliedern sich die Machine-Learning-Ansätze in drei Gruppen: Supervised Learning, Unsupervised Learning und Semi-Supervised Learning. Werfen wir einen genaueren Blick auf diese Gruppen:

3.2.1 Supervised Learning: Daten sind gelabelt

Im Fall von Supervised Learning liegt für alle Datenpunkte ein Wert für die Zielgröße vor, das sogenannte Label. **Es ist also vorher bekannt, was das trainierte Modell am Ende für jeden Datenpunkt ausgeben sollte.** Das vereinfacht die Festlegung des Trainingsziels enorm: Meistens besteht es darin, den Unterschied zwischen der Vorhersage und dem Label zu minimieren.

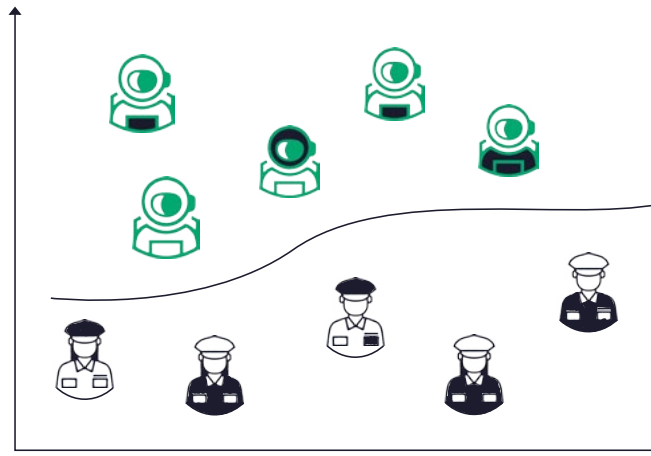


Machine-Learning-Ansätze und für sie typische Anwendungsfälle

Wenn der Algorithmus aus einer Sammlung von gelabelten Bildern, zum Beispiel, wie zuvor beschrieben, „Astronaut:innen“ identifizieren soll, dann handelt es sich also grundsätzlich um Supervised Learning.

Typische Anwendungsszenarien von Supervised Learning bilden Klassifikations- und Regressionsaufgaben, d.h. die Zuordnung von Kategorien bzw. numerischen Werten zu jedem Datenpunkt. Obwohl dies die größte Gruppe der Machine-Learning-Ansätze darstellt, ist es auch die Gruppe mit dem größten Arbeitsaufwand, denn die Bestimmung der Label für den gesamten Datensatz ist oft mit großem manuellen Aufwand verbunden.

Der Gruppe der Supervised-Learning-Ansätze stehen die Verfahren des Unsupervised Learnings gegenüber.



Klassifikation von Astronaut:innen und Polizist:innen (Supervised Learning)

3.2.2 Unsupervised Learning: Daten sind nicht gelabelt

Im Unsupervised Learning lernt der Algorithmus ohne Überwachung, Muster und Zusammenhänge in Daten explorativ zu erkennen. Die Eingangsdaten sind hier nicht gelabelt, das heißt, im Gegensatz zum Supervised Learning sind die **Werte für die Zielgröße nicht vorgegeben oder die Zielgröße selbst ist unbekannt**. Je nach Anwendung ist es ungleich schwerer, ein geeignetes Trainingsziel für den Machine-Learning-Algorithmus zu formulieren. Im Falle eines Clusterings, also der Einteilung der Datenpunkte in verschiedene Gruppen,

kann das Trainingsziel z. B. sein, dass die Gruppen untereinander möglichst verschieden, die Daten innerhalb einer Gruppe aber möglichst ähnlich sein sollen.

Es existieren unterschiedliche Typen dieses Ansatzes. Das **Clustering** teilt Daten nach Ähnlichkeiten in Kategorien ein. So lassen sich zum Beispiel **Kundengruppen** identifizieren: Kundinnen und Kunden innerhalb einer Gruppe sollen sich ähneln, sich aber gleichzeitig von denen in den anderen Gruppen stark unterscheiden. Hier ist es unmöglich, im Voraus zu wissen, was die „richtigen“ Kundengruppen sind. Das Modell soll diese spezifizieren und anschließend die Güte dieser Gruppen bewerten.

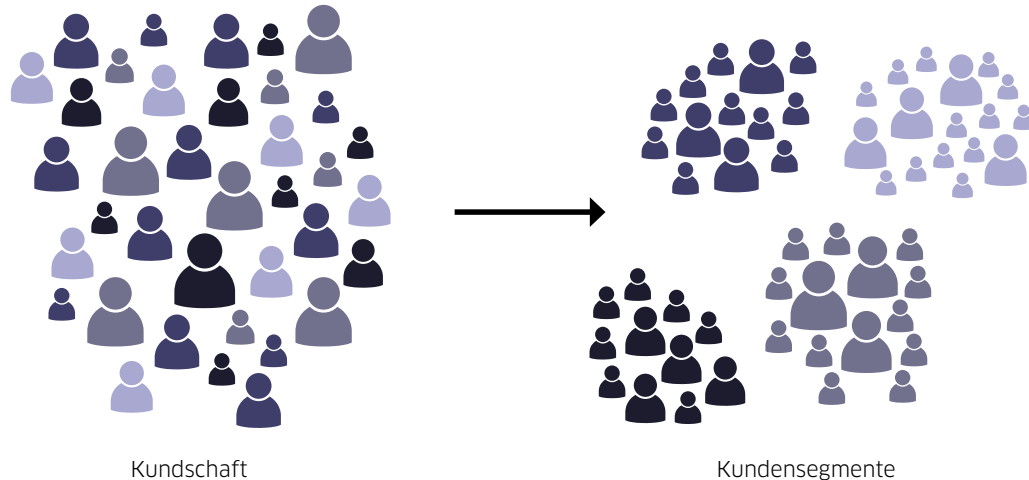
Im Rahmen der unsupervised **Anomalieerkennung** lassen sich Datenausreißer, z.B. in Produktionsprozessen oder in anderen Abläufen, frühzeitig erkennen. Sie kommt bei der Qualitätssicherung in der Produktion oder zur frühzeitigen Erkennung neuartiger Hackingangriffe zum Einsatz.

Tipp:

Was machen,
wenn es doch
mal zu viel wird?

Dimensionsreduktion ist nützlich, um Tausende Inputvariablen auf ein paar Dutzend herunterzubrechen, ohne dabei erklärende Zusammenhänge und wertvolle Informationen zu verlieren. Sie wird oft in der Vorverarbeitung für andere Modelle, vor allem neuronale Netze, genutzt.

Einen sehr großen Bereich stellt die **Analyse und Optimierung von Netzwerken** dar, beispielsweise wenn Gruppen innerhalb einer Gesellschaft identifiziert werden sollen.



Clustering von Kundensegmenten (Unsupervised Learning)

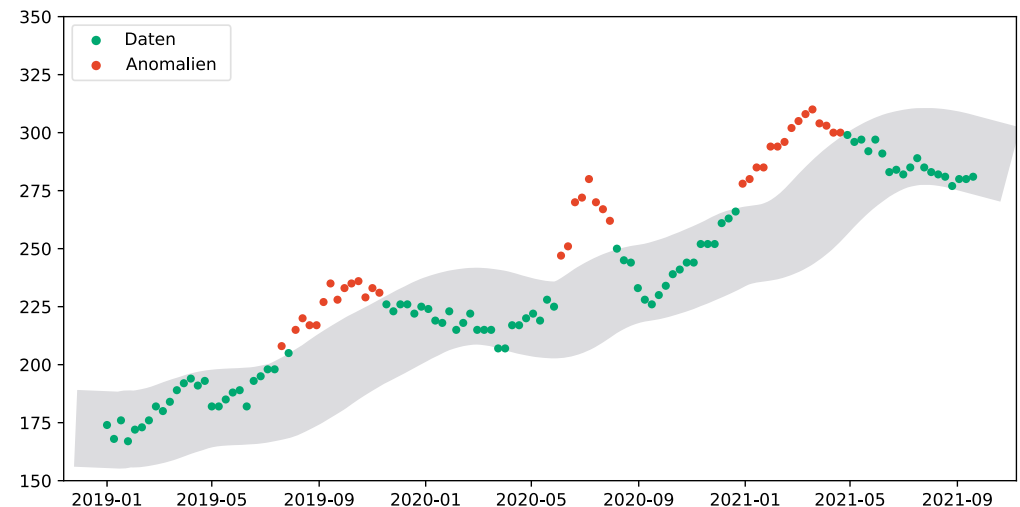
Empfehlungssysteme (Recommender-Systeme) sind den meisten von Netflix, Spotify und Co. bekannt. Das Modell gibt hier Empfehlungen zu weiteren Angeboten ab, die den bereits konsumierten ähneln.

Frequent Pattern Analysis findet oft gemeinsam auftretende Konstellationen. Das Paradebeispiel hierfür ist die Warenkorbanalyse. Die Kundschaft erhält hier Vorschläge zu Produkten, die oft gemeinsam mit dem ausgewählten Produkt gekauft werden.

Zwischen dem Supervised und dem Unsupervised Ansatz existiert noch das Semi-Supervised Learning.

3.2.3 Semi-Supervised Learning: Daten sind teilweise gelabelt

Das Semi-Supervised Learning kombiniert Supervised- und Unsupervised-Learning-Ansätze. In diesem Fall liegt nur für einen kleinen Teil der Trainingsdaten ein Wert für die Zielgröße vor. Für jeden Teil der Trainingsdaten, ob mit oder ohne Zielgröße, wird üblicherweise ein entsprechender Machine-Learning-Ansatz, Supervised oder Unsupervised, gewählt.



Anwendungsfall Anomaliendetektion auf einer Zeitreihe

Das Trainingsziel kann dann ein einzelnes kombiniertes Trainingsziel sein oder es werden für jeden Teil einzelne Trainingsziele definiert, welche gleichzeitig erfüllt werden müssen.

Unter Semi-Supervised Learning fallen überwiegend Anwendungsfälle, die theoretisch mit Supervised-Learning-Ansätzen zu lösen sind, bei denen es aber praktisch unmöglich ist, den manuellen Arbeitsaufwand zu erbringen, um alle Datenpunkte zu labeln.

Ein Beispiel: Ein Modell soll trainiert werden, um eine Emotion anhand eines geschriebenen Textes vorherzusagen. Dafür können nicht alle möglichen Konstellationen definiert und die passende Emotion vorgeschrieben werden. Stattdessen werden einige Tausend (oder mehr) Texte bewertet, auf deren Basis das Modell dann Emotionen generalisieren kann.

Exkurs I

Erfolgreiche Accso KI-Mission: Recommendation-System für ein Onlineportal

Die Aufgabe in dieser KI-Mission umfasste die Entwicklung eines Empfehlungssystems für Videobeiträge in einem Onlineportal mit dem Ziel, die Verweildauer auf der Plattform zu steigern. Dabei sollten für die Nutzende präzise, individuelle Empfehlungen eingeführt werden. Bis dahin wurden die Empfehlungen aus statistischen Rankings (z. B. beliebteste zuerst) und kuratierten Empfehlungen (d.h. manuell durch eine Redaktion ausgewählt) abgeleitet.

Warum fiel die Entscheidung auf eine ML-Lösung?

Individuelle Empfehlungen, wie sie von diversen Streaminganbietern bekannt sind, können nur mittels ML-Lösungen generiert werden. Da diese einen immensen Einfluss auf die Verweildauer der Nutzenden haben, wollte das Onlineportal diesen Vorteil nicht verschenken.

Modellaufbau mit Reinforcement Learning

Damit der Algorithmus angelernt werden und zukünftig Muster und Zusammenhänge ableiten kann, stand zunächst das Zusammentragen und die Analyse der vorhandenen Daten im Mittelpunkt. Als Grundlage dienten sämtliche Metadaten zu Videobeiträgen, z. B. Beitragstitel, Beschreibung des Beitrags, Tags sowie Daten zum Nutzungsverhalten, z. B. allgemeine Trackingdaten



Autor: Valentin Kuhn

Verfahren:
Contextual Multi-Armed Bandit

Anwendungsfall:
Bereitstellung attraktiver Videobeiträge in einem Onlineportal

Wirtschaftliches Ziel:
Erhöhung des Sehvolumens und der Nutzungsdauer des Portals, Verbesserung der User Experience und Diversifikation der angesehenen Inhalte durch Anbieten unbekannter Inhalte

Zielgröße des Modells:
IDs der Videobeiträge

Trainingssatzgröße:
mehrere Millionen Interaktionen

Anzahl Features:
2 bis 5

wie die Interaktionsdauer mit bisher empfohlenen Inhalten, angesehenen Inhalten, Reihenfolge der (letzten) Interaktionen usw.

Mit diesen ungelabelten Daten wurde der Algorithmus aufgesetzt. Ziel war es, passende Inhaltsempfehlungen zu generieren, welche die Nutzenden zum Weiterschauen anregen und somit die Verweildauer erhöhen. Das Modell wird anhand von Analysen durch ML-Ingenieurinnen und ML-Ingenieure kontinuierlich evaluiert und optimiert.

Fazit

Die bestehende kurative Empfehlungslogik wurde um die ML-Algorithmen erweitert. Dazu wurde die gesamte Empfehlungslogik umgebaut und zentralisiert. Im Ergebnis kann das Empfehlungssystem nun insgesamt schneller angepasst und die Logik deutlich besser nachvollzogen werden als bisher. Auch die Güte der kuratierten Empfehlungen kann nun mit den gleichen Tools analysiert werden wie die der automatisierten Empfehlungen. Die ML-basierten Empfehlungen haben sich inzwischen sehr gut etabliert. Aktuell wird daran gearbeitet, bei der Empfehlungslogik von manuell definierten Kategorien auf „natürliche“ Kategorien überzugehen. Diese sollen durch Clustering-Algorithmen bestimmt werden, ähnlich wie bei personalisierten Playlists von Spotify.

Herausforderungen

Auch wenn es einige Mitarbeitende im Unternehmen gab, die offen gegenüber einer ML-Lösung waren oder sogar bereits Kenntnisse mitbrachten, war ein erheblicher Überzeugungs-aufwand im Laufe des Projekts notwendig. Es wurden vor allem Bedenken zur Nachvollziehbarkeit der ML-Algorithmen oder die Angst vor Kontrollverlust der Kuratierenden geäußert. Durch die aktive Einbindung aller Beteiligten und die transparente Projektumsetzung konnten die Bedenken mit der Zeit in Vertrauen umgewandelt werden, sodass schlussendlich alle Beteiligten das Projekt unterstützen und voranbringen.



Exkurs II

Erfolgreiche Accso KI-Mission:
Proof of Concept (PoC) zur
Vorausfüllung von Buchungen
in einer Buchhaltungssoftware

In dieser KI-Mission entwickelten wir einen Proof of Concept (PoC) für ein Unternehmen, das Buchhaltungssoftware an andere Unternehmen verkauft. In einigen Fällen übersteigen die Kosten für die Bearbeitung einer Rechnung den eigentlichen Rechnungsbetrag. Die zu entwickelnde kundenspezifische ML-Lösung sollte künftig Teile eines Formulars automatisch vorausfüllen, um damit die Bearbeitungskosten für Rechnungen über Kleinstbeträge zu senken.

Warum fiel die Wahl auf einen ML-Ansatz?

Die Zielgrößen unterliegen einer hohen Komplexität. Theoretisch wäre es möglich, das Problem mit einer regelbasierten Programmierung zu lösen. In der Praxis fällt es jedoch schwer, einen Überblick über die zugrunde liegende Entscheidungslogik zu behalten. Zudem handelt es sich bei der Prüfung von Rechnungen über Kleinstbeträge um eine monotone Aufgabe. Durch Automatisierung können die Bearbeitenden viel Zeit einsparen. Auch das vom Vorstand ausgerufene strategische Ziel, ML-Wissen im Unternehmen aufzubauen und ML-Anwendungen einzusetzen, sprach für einen Einsatz von ML. Dafür bot dieses Projekt den richtigen Rahmen.



Autorin: Dr. Xenija Neufeld

Verfahren:

Decision Trees, Random Forest, Multi-Class Random Forest, Extreme Gradient Boosting (XGBoost)

Anwendungsfall:

Automatisierte Zuweisung von Rechnungen auf Buchungskonten anhand der in den Rechnungen verfügbaren Informationen

Wirtschaftliches Ziel:

Verkürzung der Bearbeitungszeit und -kosten, die durch manuelle Bearbeitung der Buchungsformulare anfällt

Zielgröße des Modells:

Anzahl richtig zugewiesener Rechnungen

Trainingssatzgröße:

Je nach Mandat zwischen 20.000 und 130.000 Datensätzen. Die Anzahl der verfügbaren Features war 20.

Anzahl Features:

Es wurde mit 3 bis 4 Features trainiert.

Modellaufbau und Herausforderungen

Die ML-Lösung sollte spezifisch für die verschiedenen Endkundinnen und Endkunden des Unternehmens entwickelt werden. Die Sondierung der Datenlage ergab, dass bei kleinerer Kundschaft nicht genug Daten vorlagen, um ein Minimum von Trainingsdaten abzudecken. Anders sah es bei den großen Kundengruppen aus. Hier lagen ausreichend Daten vor, um die erforderlichen Kriterien und auch Teilmengen pro Kundin oder Kunde zu erfüllen.

Im Laufe der Entwicklung stellte sich heraus, dass nur eine kleine Anzahl der Features nützlich war (< 5) und die Zielgrößen eine hohe Komplexität aufwiesen, d. h. konkret Dutzende bis Hunderte mögliche Werte pro Kategorie annehmen konnten. Das Modell konnte also nur einen Teil der ursprünglich avisierten Zielgrößen verlässlich vorhersagen. In gewissen Szenarien gelang dies sehr gut, z. B. wenn es sich bei der Rechnung um einen relativ kleinen Betrag handelte.

Trotz mäßiger Datengrundlage trainierten wir den Prototyp, indem wir Vorhersagen als Vorschläge markierten. Später übertrugen wir die Ergebnisse in die automatische Ausfüllung. Die Bearbeitenden profitieren bereits von einer Zeitersparnis, da die relevanten Felder des Formulars vorausgefüllt werden.

Fazit

Der entwickelte Prototyp zur Vorausfüllung von Formularen in der Buchhaltungssoftware brachte trotz mäßiger Datenlage gute Ergebnisse. Die in den abgedeckten Szenarien vorgezeigte Leistung war akzeptabel. Der Prototyp führte bereits zu guten Resultaten, erbrachte aber noch keine vollständige Vereinfachung der Prozesse für die Nutzenden und die Kundschaft. Aktuell arbeitet das Team weiter an der Datenbeschaffung sowie an der Optimierung der Modelle.

Konnte der richtige Machine-Learning-Ansatz für den vorliegenden Use Case identifiziert werden? Dann beleuchten wir jetzt den Ablauf eines ML-Projektes.

04

Takeoff, Kurskorrektur und Satellitennetz: Von der ersten ML-Lösung bis zum Betrieb der Anwendung

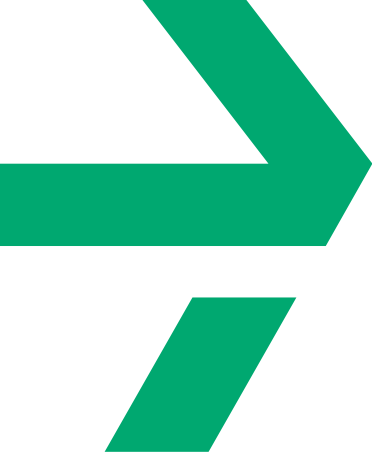
„Jeder sieht sich gerne den spektakulären Start einer neuen Mission an. Nachdem die mächtige Rakete ihre Arbeit getan hat, beendet sich der Satellit auf seinem Weg und das Rampenlicht der Medien verblasst. Dann beginnen jedoch einige der härtesten und kritischsten Arbeiten der ESA.“

Was für eine ESA-Raumfahrtmission gilt, beobachten wir auch immer wieder bei unseren Machine-Learning-Projekten. Der Auszug aus dem ESA-Blog verdeutlicht sehr gut den Aufmerksamkeitsverlauf, den auch ein ML-Projekt nimmt: Alle Seiten beobachten die Entwicklung der ersten ML-Lösung und verfolgen sie genauestens, dem Betrieb schenken sie dann nicht mehr so viel Beachtung.

Dabei sichert dieser den dauerhaften Erfolg der ML-Anwendung im Unternehmen. Daher gehen wir in diesem Whitepaper nicht nur auf die Entwicklung und die Kosten für die Erstellung einer ersten Lösung ein, sondern darüber hinaus auch auf den Betrieb einer solchen Anwendung und die damit verbundenen Kosten.

4.1 Entwicklung einer Machine-Learning-Lösung

Das Prozessmodell CRISP-DM ist unserer Meinung nach ein sehr gutes Vorgehensmodell, mit dem sich Lösungen entwickeln lassen. Es entstand bereits 1996 im Zusammenhang mit Data-Mining-Projekten und ist sehr gut auch für Machine-Learning-Vorhaben nutzbar. In unseren ML-Projekten hat es sich bewährt. Vor dem Projektstart sollten die Verantwortlichen unserer Erfahrung nach einige Voraussetzungen schaffen und Rahmenbedingungen abfragen, da diese ansonsten im CRISP-DM-Modell keine Beachtung finden.



Datenqualität und -quantität

Daten, Daten, Daten – diese sind Dreh- und Angelpunkt jeder Anwendung. Je besser die Qualität und je höher die Anzahl, desto besser das spätere Ergebnis, oder? Richtig – in der Praxis müssen wir jedoch bei der vorliegenden Datenbasis meist einen Kompromiss eingehen. So kann es durchaus vorkommen, dass das Vorhaben zunächst mit unstrukturierten Daten startet und diese erst einmal aufbereitet werden müssen. Grundsätzlich gilt aber immer, die Daten sollten verlässlich sein und keine Fehlmessungen oder Ähnliches enthalten.

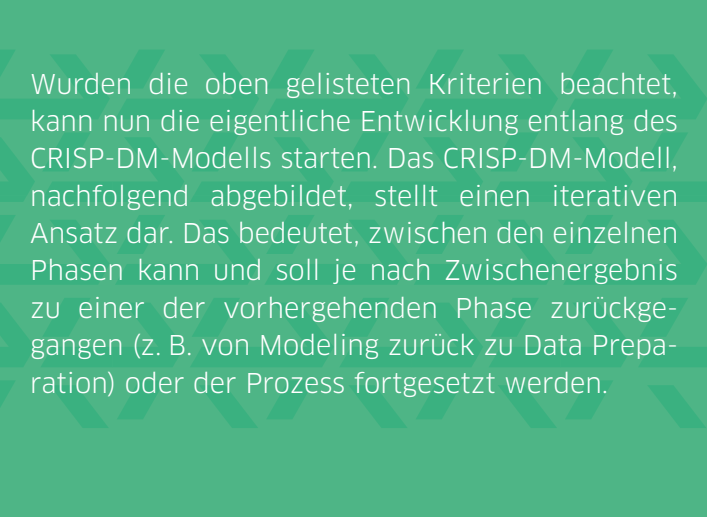
Die benötigte Datenanzahl hängt immer vom späteren Einsatzzweck und dem gewählten ML-Modell ab. Sollen die Daten beispielsweise klassifiziert werden, während einzelne Klassen aber kaum vorliegen (z. B. Patient mit einer seltenen Blutgruppe), ist zu überlegen, ob dies ignoriert werden soll oder erst mehr Daten gesammelt werden müssen. Neuronale Netze benötigen beispielsweise ein Vielfaches mehr an Daten als andere übliche ML-Modellarten. Allgemein gesprochen hat sich in unseren Projekten eine vierstellige Anzahl an Datenpunkten als ein Minimum herauskristallisiert. Unterhalb dieser Grenze sind kaum nützliche Ergebnisse zu erwarten. Für den Einsatz in der Produktion ist eine höhere Anforderung an das Vertrauen in das Modell nötig. Laut einer Studie von Dimensional Research werden deshalb in über 70 % der Projekte ein Minimum von 100.000 Datenpunkten gesetzt.

Teamzusammensetzung

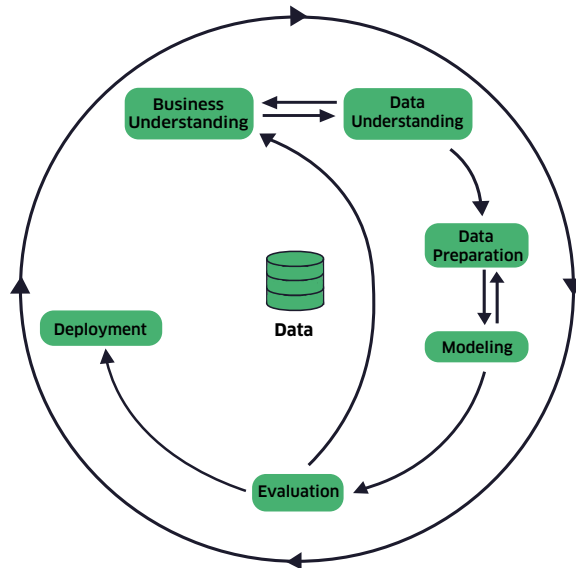
Zusätzlich zu den Daten gilt es vor Projektbeginn die relevanten Kontaktpersonen, wie z. B. Hauptverantwortliche für benötigte Datenquellen, zu bestimmen. Dedizierte Domänenexpertinnen und -experten sollten den Entwicklerinnen und Entwicklern im Projekt helfen, den Kontext der Daten und die Daten selbst zu verstehen. Gerade hierdurch lassen sich bereits im Vorfeld sehr viele Annahmen validieren.

Rechenleistung

Die geforderte Rechenleistung kann für komplexe Machine-Learning-Lösungen ein kritischer Erfolgsfaktor werden. Dies gilt vor allem, wenn neuronale Netze zur Anwendung kommen, z. B. zur Bilderkennung oder Audioerkennung. Diese benötigen eine sehr hohe Rechenleistung.



Wurden die oben gelisteten Kriterien beachtet, kann nun die eigentliche Entwicklung entlang des CRISP-DM-Modells starten. Das CRISP-DM-Modell, nachfolgend abgebildet, stellt einen iterativen Ansatz dar. Das bedeutet, zwischen den einzelnen Phasen kann und soll je nach Zwischenergebnis zu einer der vorhergehenden Phase zurückgegangen (z. B. von Modeling zurück zu Data Preparation) oder der Prozess fortgesetzt werden.



Cross-Industry Standard Process for Data Mining (CRISP-DM)

Schritt 1: Business Understanding

In dieser Phase befasst sich das Team vor allem mit dem fachlichen Anwendungsfall. Folgende Fragestellungen sind hier führend:

- Welches konkrete Problem soll gelöst werden?
- In welchem fachlichen Kontext ist die Lösung angesiedelt?
- Wie ließe sich das Problem konventionell lösen?
- Nach welcher Metrik soll die Leistung einer ML-Lösung beurteilt werden?

- Welche Key Performance Indicators (KPIs) sollen später in der Produktion zur Überwachung eingesetzt werden (z. B. die durchschnittliche Interaktionsdauer der Nutzenden)?

Schritt 2: Data Understanding

Als Voraussetzung, überhaupt in das Projekt zu starten, haben wir uns bereits zu Beginn des Kapitels mit der Datenqualität und -quantität beschäftigt. In diesem Schritt betrachten wir diese sehr viel detaillierter. Folgende Fragestellungen stehen hier im Zentrum der Überlegungen:

- Welche Datenquellen stehen zur Verfügung?
- Welche davon sind relevant und in welchem Ausmaß verfügbar?
- Was bedeuten die einzelnen Variablen?
- Wie sind die einzelnen Variablen verteilt?
- Welche Variablen erzeugen Ausreißer? Handelt es sich hierbei eventuell um Fehler?
- Was bedeuten fehlende Werte?
- Welche Beziehung besteht bei den Variablen untereinander (statistische und grafische Auswertung)?

Schritt 3: Data Preparation

Die fachlichen Fragen sind geklärt und die Daten genauestens unter die Lupe genommen. In diesem Prozessschritt wird nun eine Pipeline für die Daten aufgebaut. Diese bearbeitet die Inputdaten automatisch so, dass sie zu einem späteren Zeitpunkt in das Modell eingespeist werden können. Mit dem Aufbau der Pipeline werden folgende Tätigkeiten notwendig:

- Behandlung fehlender Werte: z. B. entsprechende Zeilen löschen vs. Spalten löschen vs. nachträglich befüllen
- Extrahieren von Teilwerten einzelner Inputvariablen: z. B. den Wochentag aus einem Datumsfeld
- Transformieren der Daten in sinnvolle Einheiten: z. B. Geldwerte, welche in Cent und Euro angegeben werden, in einen Eurowert umwandeln
- Standardisieren von Variablen: z. B. zwischen 0 und 1

Schritt 4: Modeling

In diesem Schritt widmen wir uns ausführlich dem ML-Modell. Abhängig von der Art des Anwendungsfalls (z. B. Regression) und anderen Faktoren (z. B. Transparenzansprüchen), bestimmen wir die passende Modellart – z. B. Entscheidungsbaum,

neuronales Netz, Support Vector Machine. Jede Modellart verfügt dabei über ihre ganz eigenen Parameter, die das konkrete Aussehen (wie Größe und Komplexität) eines Modells bestimmen. Auch diese Parameter haben einen Einfluss darauf, wie gut das Modell Ergebnisse erzielen kann. Deswegen geht es im Modeling-Schritt auch darum, die optimalen Werte für die Modellparameter zu finden.

Schritt 5: Evaluation

Nach der Wahl passender Modellarten werden diese jetzt anhand der im Business Understanding gewählten Metrik miteinander verglichen. In welchen Szenarien sind die Vorhersagen eines Modells gut? Welche deckt es hingegen schlecht ab? Liefern die Ergebnisse das, was aus fachlicher Sicht zu erwarten war? Hat man die Einflüsse der einzelnen Variablen richtig eingeschätzt? Ergibt die Evaluation des Modells ein zufriedenstellendes Ergebnis, so steht der Inbetriebnahme der ML-Lösung nichts mehr im Wege.

Schritt 6: Deployment

Ein entscheidender Schritt ist geschafft: Die erarbeitete Lösung wird in Betrieb genommen, um sie regelmäßig auf neue Daten anzuwenden. An dieser Stelle benötigt es noch keine Continuous Integration oder Continuous Development Pipeline. Die Mitarbeitenden beobachten zunächst, wie sich die Lösung mit den neuen Daten verhält, und vergleichen sie mit bestehenden Anwendungen.

Liefert sie auch mit neuen Daten die gewünschten Ergebnisse?

Das CRISP-Modell bietet für die Entwicklung einen sehr guten Leitfaden. Je nach Projektsituation kann und sollte man jedoch auch von den genannten Schritten abweichen. So lässt sich ein Clusteringmodell zur Fragestellung „Welche Kundengruppe habe ich?“ nicht so einfach automatisch optimieren und evaluieren wie z. B. ein Klassifikationsmodell mit der Fragestellung „Zu welcher der vordefinierten Kundengruppen gehört eine Kundin oder ein Kunde?“. Dies muss im Einzelfall entschieden werden.

Exkurs III

Was kostet die Entwicklung?

Natürlich können wir an dieser Stelle keine Pauschalkosten für die Entwicklung einer Lösung nennen. Wir versuchen jedoch, auf die Faktoren, welche die Investitionssumme beeinflussen, genauer einzugehen.

Vorhandene Expertise im Unternehmen

Welche Fachkenntnisse erfordert das Projekt? Welche davon sind bereits im Unternehmen vorhanden? Diese beiden Fragen beeinflussen in sehr hohem Maße die Gesamtkosten des Projekts. Bei vorhandenem Know-how muss geprüft werden, inwieweit die Beteiligten die tatsächlich notwendigen Kenntnisse mitbringen. Haben die Fachleute für BI und Datenanalysen beispielsweise die nötigen Kenntnisse der Kate-

gorie „Data Understanding & Exploration“? Datenbankexpertinnen und -experten können sehr gut die Aufgaben in der Phase „Data Engineering“ übernehmen. Oftmals fehlt es jedoch im Unternehmen an notwendigem Know-how in Sachen Machine Learning. Dann ist es sinnvoll, externe Fachleute wie unsere Accsonauten hinzuzuziehen, die bereits bei der Entwicklung unterstützen können und so für einen ersten Wissenstransfer in das Unternehmen sorgen.

In manchen Anwendungsfällen, z. B. beim Umgang mit Text oder Bildern, ist es sehr aufwendig, die Daten vorzubereiten. Da Teile davon, z. B. manuelles Labeln von Bildern, keine anspruchsvolle Aufgabe darstellen, lagern mehr als ein Drittel der Unternehmen solche Tätigkeiten aus (siehe Studie von Dimensional Research).



Autor: Julian Cornea

Projektlaufzeit

Neben der benötigten Expertise ist vor allem die Projektlaufzeit ausschlaggebend für die Gesamtkosten des Projekts. Die Kosten für die Entwicklung verteilen sich unserer Erfahrung nach folgendermaßen auf die einzelnen Phasen des CRISP-DM-Vorgehensmodells:

20 % Business Understanding

50 % Data Understanding und Data Preparation

20 % Modeling

10 % sonstige Aktivitäten

(z. B. Code-Verbesserung)

Je nach Erfahrungsgrad der Projektmitglieder und abhängig vom nötigen Aufwand im konkreten Anwendungsfall kann es zwei bis vier Wochen dauern, bis das Projekt evaluiert ist. Bei einem komplexeren Problem (z. B. Bilderkennung) kann sich dies auf mehrere Monate erstrecken. Laut der Studie von Dimensional Research müssen 90 % der Unternehmen unerwartete Schwierigkeiten einkalkulieren, wenn sie Daten vorbereiten.

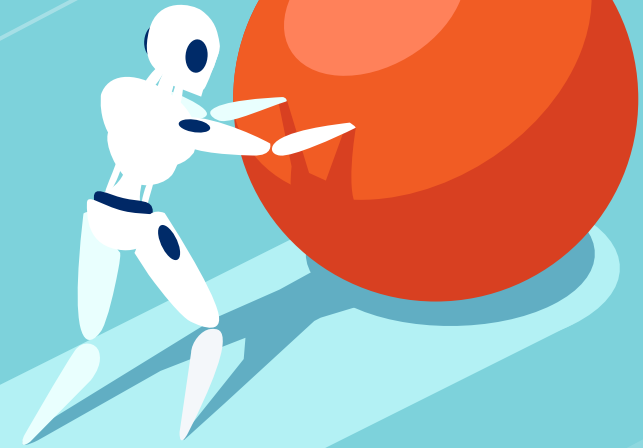
Optional:

Kosten für „mehr“ Daten und Rechenkapazität

Je nach Projekt ist es eventuell erforderlich, neue Daten für die Anwendung zu beschaffen. Je nach Datenbasis kann das teuer werden - gerade wenn es sehr viel manuellen Aufwand erfordert, die Daten bereitzustellen. Trotz der Kostenoptimierung durch Outsourcing können Preise durchaus

im vier- oder sogar fünfstelligen Bereich liegen. Dies gilt vor allem für Use Cases in den Bereichen Bild- und Videoerkennung. Bei Textverarbeitung fallen die Preise tendenziell noch höher aus, da das umsetzende Team gewisse sprachliche Voraussetzungen erfüllen muss.

Rechenintensive neuronale Netze benötigen außerdem sehr viel Rechenleistung, die eventuell in der Cloud gemietet werden muss. Auch hier können die Kosten für die Entwicklungsphase je nach Bedarf im drei- bis fünfstelligen Bereich liegen.



4.2 Betrieb und Weiterentwicklung einer Machine-Learning-Lösung

Sobald ein Modell produktiv genutzt werden soll, zerfällt der lineare Aufbau des CRISP-DM-Vorgehensmodells. Ab sofort erfolgt die Weiterentwicklung nur noch an den Stellen, an denen Bedarf besteht. So stellt sich bei einem Modell z. B. heraus, dass Ausreißer nicht so gut abgefangen werden wie erwartet. In diesem Fall ließe sich das Modell mit anderen Parametern neu trainieren, ohne dass man sich nochmals mit der Datenpipeline beschäftigen müsste. In einem anderen Fall ändert sich die Formatierung einer Datenquelle. Hier muss nur die Data Preparation Pipeline angepasst werden, jedoch nicht das eigentliche Modell. Vielleicht ist es auch sinnvoll, einen A/B-Test durchzuführen, um das Modell mit der etablierten Lösung zu vergleichen. Dies wirkt sich weder auf die Datenpipeline noch auf das Modell aus und erfordert kein neues Deployment.

Neue Technologien und benötigtes Fachwissen

Bei einer produktiven ML-Anwendung ergeben sich ganz neue Fragestellungen und Herausforderungen, die zusätzliche Technologien und Fachwissen erfordern. Beispielsweise soll die Vorhersage des Modells in eine andere Anwendung einbezogen werden. Hierfür ist eine Schnittstelle unerlässlich. Die Vorhersagen werden nun in Echtzeit benötigt. Schwankt die Auslastung des Systems dabei sehr?

Dann muss die Anwendung skalierbar gestaltet werden, was sich wiederum durch Architekturmaßnahmen wie Lastverteilung oder die Nutzung autoskalierender Services erzielen lässt.

Eine verantwortliche Person sollte im Betrieb Änderungen in den Datenquellen verfolgen. Damit wird gewährleistet, dass man z. B. frühzeitig erkennt, wenn Änderungen in der Data Preparation Pipeline notwendig werden. Dies ermöglicht eine reibungslose Nutzung der Anwendung. Nicht nur die Datenquellen sollten stetig im Auge behalten werden, sondern auch die Modellleistung. Ein automatisches Monitoring erkennt Leistungsabnahmen und kann entsprechend einen Alarm auslösen.

Notfalls kann man auf eine vorherige Modellversion zurückschwenken. Dies erfordert ebenfalls eine CI/CD-Pipeline und eine vorhandene Monitoring-Infrastruktur.

Vorhandene Expertise im Unternehmen

Nun ist auch ein guter Zeitpunkt, darüber nachzudenken, ML-Engineering-Kenntnisse im Unternehmen aufzubauen. Dies kann über die Einstellung eines erfahrenen ML-Engineers, die Weiterbildung vorhandener Mitarbeitender oder durch die Unterstützung eines Dienstleisters wie Accso erfolgen. Fällt die Entscheidung auf die interne Weiterbildung, eignen sich besonders diejenigen mit tiefen Statistikkenntnissen, da sich die zugrunde liegende Mathematik stark

überschneidet. Alle weiteren notwendigen Kenntnisse für den Betrieb sind wahrscheinlich bereits im Unternehmen vorhanden. Der einzige Bereich für den womöglich kein Experte vorhanden ist, ist das Deployment eines ML-Modells.

Gut zu wissen:

Ist die Infrastruktur für ML-Anwendungen bereits im Unternehmen vorhanden?

In den ersten Phasen einer produktiven ML-Anwendung können hohe Kosten entstehen, wenn die Infrastruktur im Unternehmen erst noch geschaffen werden muss, wie z. B. eine Trainingsumgebung, eine CI/CD-Pipeline oder ein Monitoringsystem.

Benötigt die Anwendung regelmäßig neue Datenquellen?

Kann die Anwendung ohne große Änderungen in den angebundenen Datenquellen verwendet werden, fällt relativ wenig Aufwand an, sofern das Unternehmen ein gutes Monitoringsystem nutzt. Müssen neue Datenquellen angebunden werden, bedeutet dies regelmäßig hohen Aufwand in den Bereichen „Data Understanding“ und „Data Preparation“.

Aus dem Accsonauten-Alltag: Erfolgsfaktoren und Lessons Learned aus der Praxis

Zu guter Letzt werfen wir einen Blick hinter die Kulissen unserer täglichen Arbeit und gehen auf die aus unserer Sicht wichtigsten Erkenntnisse aus verschiedenen KI-Projekten ein. Welche Faktoren sind für den Erfolg eines ML-Projektes zu berücksichtigen? Auf was sollte man unbedingt achten? Was hingegen sollte man tunlichst vermeiden?

5.1 Erfolgsfaktoren für ein KI-Projekt

„Wir fliegen, wir fliegen“, sollen die ersten Worte gewesen sein, die Juri Alexejewitsch Gagarin am 12. April 1961 als erster Mensch im Weltraum von sich gegeben hat. Damit diese Unternehmung überhaupt möglich war, mussten viele Einzelschritte koordiniert, abgestimmt und verbessert werden. Da wir bereits zahlreiche Machine-Learning-Anwendungen realisiert haben, sind gewisse Faktoren in den Vordergrund gerückt, die zum Gelingen eines ML-Projektes beitragen. Diese lassen sich zum einen in technische und organisatorische Rahmenbedingungen untergliedern. Zum anderen

gehen wir darüber hinaus im Speziellen auf die Produktion ein, da dieser Bereich eine große Herausforderung für viele Unternehmen darstellt.

5.1.1 Technische Aspekte

In technischer Hinsicht müssen sich Unternehmen vor allem mit der Bereitstellung der benötigten Rechenleistung und der Integration in die bestehende Infrastruktur beschäftigen. Gerade neuronale Netze erfordern eine hohe Rechenleistung, vor allem sehr leistungsstarke Prozessoren, wie die von Grafikkarten oder speziellen TPUs, um ML-Modelle zu trainieren. Hierbei ist besonders darauf zu achten, ob die relevanten neuronalen Netze oder notwendige Entwicklungsframeworks unterstützt werden (z. B. NVIDIA Cuda). Die benötigte Rechenleistung kann entsprechend in physischer Form gekauft oder auch in der Cloud gemietet werden. Klassische ML-Modelle, z. B. Entscheidungsbäume, kommen hingegen auch mit weniger Rechenleistung aus.

5.1.2 Organisatorische Faktoren

Transparenz in Methodik und Modellwahl

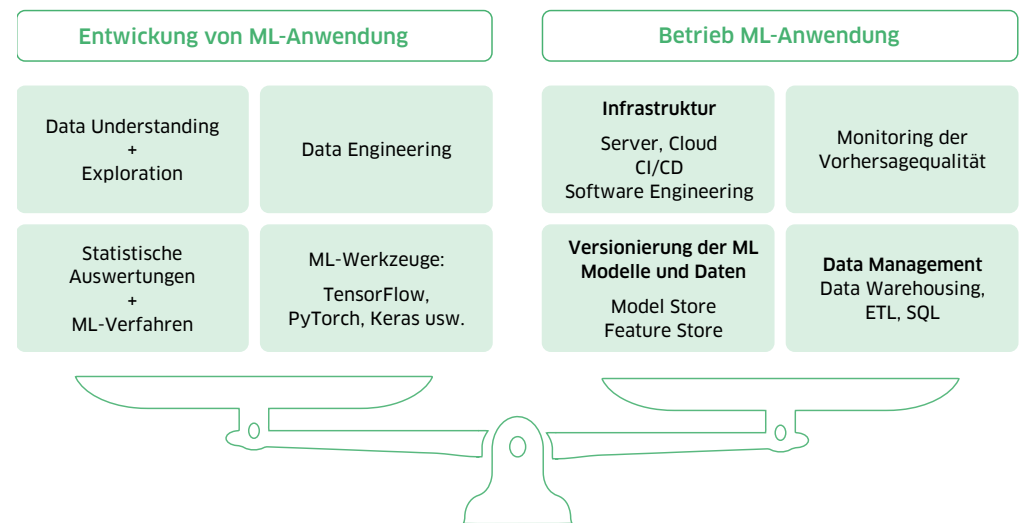
Wie bei den meisten IT-Projekten bedarf es auch bei Machine-Learning-Unternehmungen der Unterstützung durch Entscheidende. Transparenz entsteht, wenn sich ML-Expertinnen und Experten und Fachleute schnell und regelmäßig inhaltlich austauschen. Dies stärkt das Vertrauen in das Projekt und sichert zusätzlich den Erfolg. Dieser Austausch sollte sich nicht nur auf die Phase „Business Understanding“ beschränken, sondern sich auch auf die Schritte „Data Understanding“ und „Data Preparation“ ausdehnen.

Transparenz gilt es auch bei der Wahl des passenden ML-Modells zu berücksichtigen. Über die Anforderungen an die Nachvollziehbarkeit eines Modells sollte man sich so früh wie möglich Gedanken machen. Oft erzielen „Black Box“-Modelle, die man schwer nachvollziehen kann, bessere Ergebnisse als nachvollziehbare Modelle. Je nach Gebiet kann aber aus rechtlichen oder organisatorischen Gründen vorgeschrieben sein, wie transparent ein KI-Modell sein muss. Ein nachvollziehbares Modell kann beispielsweise Fragen beantworten wie: Warum sind gewisse Variablen wichtig bzw. unwichtig? Wie stark beeinflusst eine Variable das Ergebnis? Hat sie eine positive oder negative Auswirkung? Welche Entscheidungslogik liegt einer gegebenen Vorhersage zugrunde?

Fachkenntnisse als Fundament

Weder die Entwicklung noch der Betrieb einer ML-Anwendung kommen ohne die notwendigen Fachkenntnisse im Unternehmen aus. Die Grafik verdeutlicht die notwendigen Skills - unterteilt in die Entwicklung und den Betrieb der Anwendung. Teile der genannten Fähigkeiten entstammen dem CRISP-DM-Modell.

Entwicklung und Betrieb einer ML Plattform



Entwicklung und Betrieb einer ML-Plattform

5.1.3 Einsatz in der Produktion

Nur etwa die Hälfte der durchgeführten ML-Projekte schaffen es in die Produktion. So lautet unter anderem ein Ergebnis der Studie von Dimensional Research. Daher gehen wir gesondert auf diesen Punkt ein. Bei dem Betrieb einer ML-Anwendung sind folgende Aufgaben für einen langfristigen Erfolg zu bewältigen:

Solides Softwareengineering:

- Einsatz einer CI/CD-Pipeline
- Monitoring Stack
- Skalierbares Setup der Infrastruktur
- Testen (manuell und automatisiert)

Monitoring der KI-Vorhersagen und KPIs:

ML-Anwendungen sind fragiler als klassische Software, denn die Qualität des Outputs kann sich sehr plötzlich verändern. Daher müssen die KI-Vorhersagen und definierten KPIs, z.B. „Wie oft klicken Nutzende auf meine Empfehlungen?“, im Blick behalten werden. Ein vorher definiertes Alarmsystem weist im besten Fall automatisch auf einen plötzlichen Qualitätsverlust hin.

Versionierung von Modellen und Daten:

Ein Qualitätsverlust ist eingetreten! Dies liegt wahrscheinlich daran, dass eine aktuelle Version

mit neuen Daten beliefert wurde, welche in dieser Form noch keine Repräsentation in den Trainingsdaten vorfinden. Jetzt ist es wichtig, auf die alte Modellversion zurückgreifen zu können. Eine Versionierung mit den relevanten Parameterwerten ist unerlässlich. Idealerweise umfasst dies auch Teile der Daten. Unternehmen können auf diese Weise besser nachvollziehen, was für die Unterschiede in der Qualität der Vorhersagen verantwortlich ist. Dies gilt vor allem für Modelle, welche in Echtzeit lernen und somit nicht sehr streng kontrolliert werden können.

A/B-Tests:

Liefert die ML-Anwendung wirklich den erhofften Mehrwert? Diese Frage können nur A/B-Tests von verschiedenen Modellen, also ML-Modelle versus herkömmliche Lösungen, beantworten. Um die Kosten für ML-Projekte zu rechtfertigen, lohnt sich eine Gegenüberstellung. Eventuell ist der Nutzen geringer als erwartet.

Frequenz des Modelltrainings:

Diese Anforderung ist vollkommen anwendungsabhängig. Reicht ein großes Training einmal pro Woche während geringerer Auslastungszeiten (Batch Training) aus oder wird ein Echtzeittraining benötigt? Je näher sich die Frequenz einem Echtzeittraining nähert, desto höher das Risiko, dass die Vorhersagen plötzlich sehr viel schlechter werden.

5.2 Dos and Don'ts bei der Umsetzung von KI-Projekten

Neben den erläuterten technischen und organisatorischen Faktoren bzw. den Rahmenbedingungen für den Betrieb einer ML-Anwendung stellen unsere Accsonauten nun weitere Best Practices im Umgang mit ML-Projekten vor.

5.2.1 Dos

Think global, act local!

Für das erste Machine-Learning-Projekt lohnt sich am besten ein Anwendungsfall, der möglichst überschaubar, aber gleichzeitig ausbaufähig ist. Mit Blick auf die internen Prozesse eignen sich vor allem aufwendige monotone Routineaufgaben, die zum Beispiel in Excel ausgeführt und bereits mit verlässlichen und relevanten Daten versorgt werden.

Messbare Mehrwerte präsentieren

Für die anfängliche Überzeugungsarbeit und die langfristige Rechtfertigung der benötigten Ressourcen sollte man sich bereits zu Beginn Gedanken über messbare Vorteile der fertigen ML-Anwendung machen. Die Business-Benefit-Matrix aus dem BI Guidebook von Rick Sherman liefert hierfür sehr gute Messgrößen, wie beispielsweise:

- Umsatzerhöhung pro extra € Vertriebsaufwand
- Reduzierung der Lagerkosten um ... %
- Anzahl der erkannten Betrugsversuche steigt um ... %
- Vermeidung einer Strafe von ... €, welche z. B. anfallen würde, wenn neue Gesetzesanforderungen nicht erfüllt werden
- Ermöglicht den Markteinstieg in weiteren Ländern
- u. v. m.

Diskriminierung minimieren und Bias vermeiden

Amazon konnte 2017 seinen Einstellungsprozess nicht automatisieren, da die eingesetzte KI Frauen diskriminierte. Das Modell orientierte sich an Merkmalen, die bei eingestellten Mitarbeitenden verbreitet waren. Wie in der IT-Branche üblich, bildeten auch bei Amazon Männer den Großteil der Belegschaft, sodass das Modell einen Bias hin zu männlichen Bewerbern erlernt hatte. Am Ende wurden weibliche Bewerber benachteiligt, wenn diese zum Beispiel an einer Universität studiert hatten oder Teil eines Vereins waren, welche ausschließlich Frauen zulassen. ML-Modelle tendieren dazu, nach Variablen zu diskriminieren, die unausgewogen verteilt sind. Es ist wichtig, sich mit Variablen zu beschäftigen, nach denen das eingesetzte Modell nicht diskriminieren darf.

So hat zum Beispiel die National Association of Insurance Commissioners im Jahr 2020 mehrere Prinzipien für die Versicherungsbranche definiert, zu denen unter anderem ein ML-Diskriminierungsverbot gegen „geschützte Gruppen“ gehört. Forschende bei IBM haben bereits im Jahr 2017 detaillierte Verfahren veröffentlicht, welche das Diskriminierungspotenzial in ML-Projekten einschränken, ohne das Erfolgspotenzial zu beeinträchtigen.

Datenteilmengen einbehalten

In der Prototypenphase sollte ein Teil der Daten einbehalten werden und nicht zum Trainieren des Modells genutzt werden. So kann unabhängig von den Entwicklenden evaluiert werden, wie gut oder schlecht ein Modell abschneidet, wenn es mit komplett neuen Daten in Kontakt kommt. Spätestens nach dem Deployment würde sich dies ohnehin offenbaren. Mit diesem Vorgehen sparen Unternehmen Zeit und Geld, falls die Ergebnisse im Evaluierungsschritt nicht zufriedenstellend sind.

Zusätzlich ist es sinnvoll, dass jeweils eine Gruppe von Data-Engineers die Daten beschafft und säubert sowie eine weitere Gruppe von Data-Scientists bzw. ML-Engineers die Daten analysiert und die Modelle trainiert. Dabei erhält die zweite Gruppe nur auf einen Teil der Daten, ca. 70 %, Zugriff.

Datenmenge schlägt Modelloptimierung

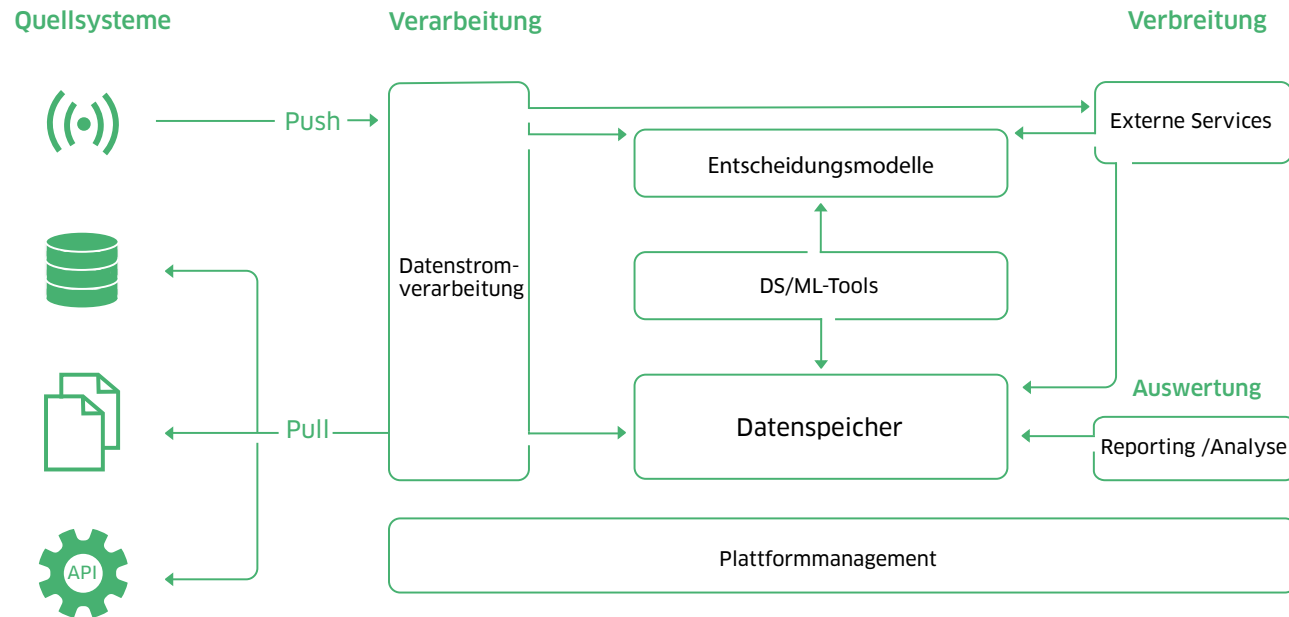
Steht das Projektteam vor der Wahl, entweder zusätzliche Datenquellen anzuzapfen bzw. neue Daten zu sammeln oder das ML-Modell weiter zu optimieren, so empfehlen wir meist die erste Option. Ein vergleichsweise „primitives“ Modell mit sehr vielen Daten liefert bessere Ergebnisse als ein hochkomplexes Modell mit wenigen Daten. Eine zunehmende Menge relevanter Daten minimiert den Unterschied zwischen den eingesetzten Modellen. Das demonstriert das berühmte Paper „Scaling to Very Very Large Corpora for Natural Language Disambiguation“ von Michele Banko und Eric Brill aus dem Jahr 2001.

Das Rad nicht neu erfinden

Wenn es für einen Anwendungsfall bereits sehr gute vortrainierte Modelle gibt, kann es sehr sinnvoll sein, auf diese zurückzugreifen. Die Nutzung vortrainierter Modelle sollte vorsichtig abgewogen werden. Grundsätzlich ist dies ein möglicher Ansatz, die Entwicklung zu beschleunigen.

Fundament für spätere Architektur gießen

Bereits in der Entwicklungsphase lohnt es, sich Gedanken über die spätere Architektur der produktiven ML-Anwendung zu machen. Die Basis für ein erfolgreiches und gut integriertes ML-Projekt ist eine Datenplattform wie nach der Referenzarchitektur von Accso.



Referenzarchitektur einer Datenplattform

Dokumentation schafft Erfolg

Damit Verantwortliche eine ML-Anwendung erfolgreich verwalten können, sollten sie die folgenden Aspekte dokumentieren:

- Welche Daten werden genutzt, um das Modell zu trainieren und zu testen?
- Welche Programme und Algorithmen werden genutzt, um das Modell zu trainieren und zu testen?
- Welche Personen evaluieren die Leistung eines Modells?

→ Wann und wo wird das Modell eingesetzt?

→ Wann und wie oft wird das Modell aktualisiert?

Dolmetschen zwischen den Welten

Eine verantwortliche Person sollte im Projekt den Output von Modellen interpretieren und an Führungspersonen ohne statistische Kenntnisse verständlich kommunizieren (z. B. mithilfe von Grafiken und nachvollziehbaren Metriken). Das schafft nicht nur Transparenz und Verständnis, sondern gewährleistet letztlich den Erfolg.

Facts and Figures

Grundsätzlich ist eine datenbasierte Unternehmenskultur zu befürworten. Auch die beste ML-Anwendung nutzt nicht viel, wenn Entscheidungen weiterhin auf Bauchgefühl statt auf Daten basieren. Liefert eine ML-Anwendung gute Ergebnisse, müssen im nächsten Schritt die betroffenen Entscheiderinnen und Entscheider überzeugt werden, diese Ergebnisse auch zu nutzen. Laut dem Harvard Business Review scheitern 95 % der KI- und Big-Data-Initiativen an organisatorischen und kulturellen Hürden und nicht, weil die Technologie zu kompliziert ist. Wenn die Unternehmensführung also genauso offen für datenbasierte Argumente ist wie Jeff Bezos, dann fallen viele Hürden weg.

Datenschutz und Privatsphäre

Mithilfe von ML können Unternehmen ihrer Kundschaft individuell zugeschnittene Angebote unterbreiten. Dabei sollten immer die Präsentation und Wirkung solcher Ergebnisse berücksichtigt werden. Facebook-Nutzerinnen und Nutzer könnten beispielsweise personalisierte Werbung über Produkte, über die sie sich kürzlich mit Bekannten ausgetauscht haben, als Eingriff in die Persönlichkeitsrechte und Privatsphäre werten. Sie fühlen sich überwacht. Die Kundschaft von Netflix und Spotify dagegen bewertet Empfehlungen oder individuell zusammengestellte Playlisten tendenziell als sehr nützlich und positiv. Eine qualitative Untersuchung der Konrad-Adenauer-Stiftung über verschiedene Bevölkerungsgruppen hinweg zeigt

deutlich die ambivalente Haltung der Menschen gegenüber KI. Alle Bevölkerungsschichten nehmen das Spannungsfeld von Sicherheit einerseits und persönlicher Privatsphäre sowie Datensicherheit andererseits ähnlich kritisch wahr. Hier ist eine entsprechend sensible Vorgehensweise erforderlich.

5.2.2 Don'ts

Begründete Zweifel – Ergebnisse hinterfragen

Der Prototyp liefert eine Genauigkeit von 95 %? Sehr gut, dann ab damit in die Produktion! Oder vielleicht doch nicht? „Sehr gute“ Metriken und sehr gute Ergebnisse eines Prototyps sollten nicht einfach so hingenommen werden. Oftmals lohnt sich ein zweiter Blick auf die wichtigsten Erfolgskriterien: War das Datenset gesplittet? Passt die Modellkomplexität zur verfügbaren Datenmenge? Reichen 1.000 Datenpunkte, um ein komplexes neuronales Netz zu trainieren?

One size fits all? – Neuronale Netze sind kein Allheilmittel

Neuronale Netze eignen sich sehr gut, um Zusammenhänge in Datenmengen mit einem sehr hohen Anteil irrelevanter Daten und einem kleinen Anteil relevanter, zusammenhängender Daten zu finden. Deshalb sind sie die Norm, wenn es z. B. um Text-, Bild- oder Audiodaten geht. Jedoch haben neuronale Netze auch große Nachteile. Unter anderem benötigen sie ein sehr hohes Datenvolumen und viel Zeit zum Trainieren. Zusätzlich sind sie stark intransparent, sodass man die Entscheidungen kaum nachvollziehen kann. Sie erfordern zudem ein deutlich detaillierteres theoretisches Wissen und weitere Technologien, wie z. B. TensorFlow oder PyTorch, inklusive spezialisierter ML-Fachkräfte, um diese Anforderungen abzudecken. Alle ML-Anwendungsfälle

mit dem „Neuronale Netze“-Hammer zu lösen, ist daher nicht der richtige Ansatz. Es gilt, spezifisch auszuwählen.

„Du kommst hier nicht rein“ – vom zu späten Einbeziehen der End User

End User können auf organisatorische und fachliche Zusammenhänge hinweisen, die bestimmte Ansätze sinnlos erscheinen lassen, z. B. Daten, welche zum Zeitpunkt der Vorhersage nicht vorhanden sind. Die Entwickelnden können die Verlässlichkeit der ML-Anwendung früh transparent halten, damit die End User verstehen, wie und warum Vorhersagen getroffen werden.

Arnold, Norbert; Frieß, Hans-Jürgen; Roose, Jochen und Werkmann, Caroline: Wenn die KI unser Assistent bleiben kann, dann können wir viel draus ziehen, Künstliche Intelligenz in Einstellungen und Nutzung bei unterschiedlichen Milieus in Deutschland, Konrad Adenauer Stiftung, 2020

Banko, Michele, und Brill, Eric: Scaling to Very Large Corpora for Natural Language Disambiguation, Proceedings of the 39th Annual Meeting of the Association for Computational Linguistics, Toulouse: ACL, 2001

Dastin, Jeffrey: Amazon scraps secret AI recruiting tool that showed bias against women, 2018, <https://www.reuters.com/article/us-amazon-com-jobs-automation-insight-idUSKCN1MK08G>

Davenport, Thomas H. Mittal, Nitin: How CEOs Can Lead a Data-Driven Culture, Harvard Business Review, 2020, <https://hbr.org/2020/03/how-ceos-can-lead-a-data-driven-culture>

Dimensional Research: Artificial Intelligence and Machine Learning Projects Are Obstructed by Data Issues, Global Survey of Data Scientists, AI Experts and Stakeholders, 2019, <https://cdn2.hubspot.net/hubfs/3971219/Survey%20Assets%201905/Dimensional%20Research%20Machine%20Learning%20PPT%20Report%20FINAL.pdf>

Mittelstand Digital: KI-Kochbuch, BSP Business School Berlin – Hochschule für Management GmbH, 2021

National Association of Insurance Commissioners: National Association of Insurance Commissioners (NAIC) Principles on Artificial Intelligence (AI), 2020, https://content.naic.org/sites/default/files/inline-files/AI%20principles%20as%20Adopted%20by%20the%20TF_0807.pdf

Sherman, Rick: Business Intelligence Guidebook: From Data Integration to Analytics, Waltham, MA (USA): Morgan Kaufmann, 2014

Varshney, Kush R.: Reducing discrimination in AI with new methodology, IBM, 2017, <https://www.ibm.com/blogs/research/2017/12/ai-reducing-discrimination/>



Ihre Ansprechpartner

für Data Science und Machine Learning
machine-learning@accso.de

Wir hoffen, wir konnten Ihnen einen guten und verständlichen Einblick in den Kosmos von Künstlicher Intelligenz und Machine Learning geben.

Wurde Ihre Lust auf erste eigene Erfahrungen geweckt?

Melden Sie sich gerne bei uns!



Dr. Xenija **Neufeld**



Valentin **Kuhn**

SHARING YOUR CHALLENGE



ACCELERATED SOLUTIONS

Accso – Accelerated Solutions GmbH

Tel. +49 (6151) 13029-0
Fax: +49 (6151) 13029-10
machine-learning@accso.de
accso.de

Hilpertstraße 12
Rahmhofstraße 2
Im Medienpark 6a
Balanstraße 55
Clock Tower, 302

| 64295 Darmstadt
| 60313 Frankfurt a. M.
| 50670 Köln
| 81541 München
| CV&A Waterfront, Cape Town 8002, ZA

