

KI UND ZFP: CHANCEN UND RISIKEN

Dr. Aschwin Gopalan
Rohmann GmbH

AI



Überblick

- KI - Was ist das? Eine kurze Geschichte
- Chancen, Risiken und Herausforderungen
- KI und andere Verfahren in der ZfP

AI

Warum reden wir in der ZfP über KI?

Weil es alle tun...

- Treibende Faktoren für KI in der ZfP
 - Automatisierungsdruck
 - Fachkräftemangel
 - Kundenerwartungen („KI kann das bestimmt besser“)
- Warum müssen wir KI wenigstens ansatzweise verstehen?
 - Um Chancen und Risiken erkennen zu können
 - Um unrealistische Erwartungen zu dämpfen



KI – WAS IST DAS?

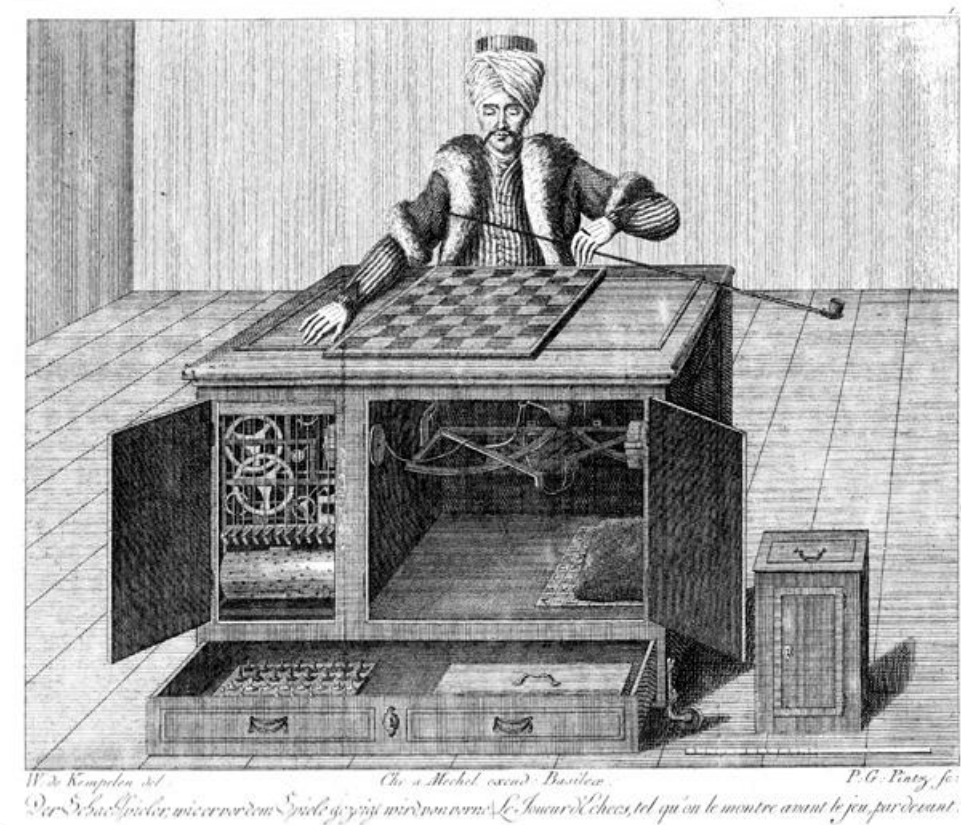
Eine kurze Geschichte und wichtige
Verfahren

AI

Künstliche Intelligenz – Ein Menschheitstraum

Eine Maschine, die denken kann!

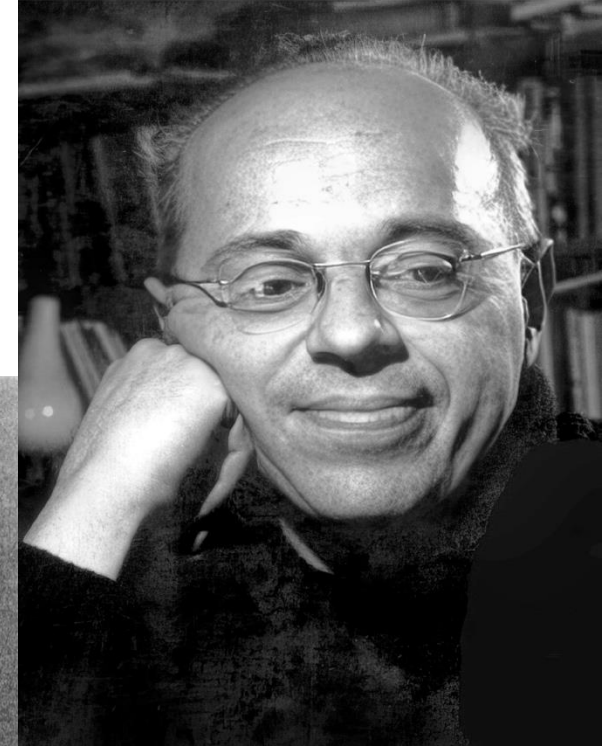
- Der Schachtürke (The Mechanical Turk)
- Erfunden von Wolfgang von Kempelen 1769
- Ist durch ganz Europa und Amerika getourt und hat gegen viele Schachmeister gewonnen
- Niemand konnte die Funktionsweise erklären
- Leider keine KI, im inneren war ein sehr guter kleinwüchsiger Schachspieler verborgen, der den „Roboter“ gesteuert hat.
- Edgar Allen Poe schrieb ein Essay mit 17 Punkten, die dafür sprachen, dass ein Mensch agiert.
- Wichtigstes Argument: Er verliert manchmal



Frühe Ideen in der Science Fiction: Lem, Asimov et. Al.

Ethische Fragen der KI, bevor es sie gab.

- Frühe Voraussagen
 - Stanislaw Lem, Summa Technologiae 1964: Deep Learning, Reinforcement Learning, selbstoptimierende Netze
- Wie kann sich der Mensch schützen?
 - Asimov: Die Robotergesetze, fest einprogrammiert in jedes positronisches Gehirn
 - Roboter und KI als Waffen? Heute Realität!
- Haben Roboter Persönlichkeitsrechte?
 - Ab wann hat eine KI Bewusstsein?
 - Erwachsen daraus Rechte?
 - „Der zweihundertjährige mit Robin Williams“



Licensed under CC-3.0
Attribution Share alike by
Wojciech Zemek

Alan Turing (1912 – 1954)

Und das “Imitation Game”

- Englischer Mathematiker
- Stellt sich eine Maschine vor, die weit über ihre ursprüngliche Programmierung hinaus lernen kann
- Die damaligen Computer waren viel zu Leistungsschwach -> Gedankenexperiment
- Stellt „The Imitation Game“ (heute: „Turing Test“) als Kriterium für eine künstliche Intelligenz, die dem Menschen ebenbürtig ist vor.
- GPT-4 hat den Test 2025 bestanden!



Die Dartmouth Konferenz

Sommer 1956

- Eingeladen von John McCarthy, Mathematikprofessor am Dartmouth College
- Gedankenaustausch der führenden Köpfe über „Denkmaschinen“
- Teilnehmer u.a. Marvin Minsky (Harvard), Nathaniel Rochester (IBM), Claude Shannon (Bell Labs)
- “Every aspect of learning or any other feature of intelligence can in principle be so precisely described that a machine can be made to simulate it”
- Oft als „Geburtsstunde der KI“ bezeichnet



Copyright the Minsky Family

ELIZA: Der erste Chatbot

Joseph Weizenbaum (MIT) 1966

- Sehr einfacher Chatbot, der einen „Therapeuten“ simulieren soll
- „Hey, I am ELIZA, tell me your problems“
- Eigentlich gedacht um zu zeigen, wie primitiv Künstliche Intelligenz ist
- Viele Menschen haben dennoch eine Beziehung zu ELIZA aufgebaut

```
Welcome to

          EEEEEEE LL      IIII  ZZZZZZZZ  AAAAA
          EE      LL      II      ZZ      AA  AA
          EEEEE  LL      II      ZZZ      AAAAAA
          EE      LL      II      ZZ      AA  AA
          EEEEE  LLLLLL  IIII  ZZZZZZZZ  AA  AA

Eliza is a mock Rogerian psychotherapist.
The original program was described by Joseph Weizenbaum in 1966.
This implementation by Norbert Landsteiner 2005.

ELIZA: Is something troubling you ?
YOU:   Men are all alike.
ELIZA: What is the connection, do you suppose ?
YOU:   They're always bugging us about something or other.
ELIZA: Can you think of a specific example ?
YOU:   Well, my boyfriend made me come here.
ELIZA: Is it important to you that your boyfriend made you come here ?
YOU:   He says I'm depressed much of the time.
ELIZA: I am sorry to hear that you are depressed.
YOU:   It's true. I am unhappy.
ELIZA: Can you explain what made you unhappy ?
YOU:   █
```


Der „KI Winter“ und Deep Blue

Mitte der Siebziger bis Mitte der Neunziger

- Forschung lief weiter, aber die Enttäuschung war groß
- Fördermittel wurden gekürzt, kaum Investitionen
- Expertensysteme und andere Entwicklungen wirkten nicht besonders intelligent und lernfähig
- Dann kam Deep Blue: Schachcomputer von IBM
- Regelbasiertes System, kein echtes Machine Learning
- Gewinnt gegen Gary Kasparov, amtierender Weltmeister!



Der Winter ist zu Ende: AlexNet

CNNs erobern die Welt

- Neuronale Netze simulieren ein wenig unser Gehirn
- Idee ist nicht neu
- ImageNet Challenge: Klassifizierung von Bildern
- Vor 2012 waren immer klassische Verfahren und Machine Learning die Gewinner
- Gewinner 2012: AlexNet (Alex Krizhevsky, Ilya Sutskever, Geoffrey Hinton)
 - Tiefes CNN, trainiert mit Grafikkarten (GPU) mit massiven Datenmengen
 - Gewinnt mit sehr großem Abstand!
- Geburtsstunde des „Deep Learning“
- Fokus von „Algorithmen“ zu „Daten“ verschoben!

ImageNet Classification with Deep Convolutional Neural Networks

By Alex Krizhevsky, Ilya Sutskever, and Geoffrey E. Hinton

Abstract

We trained a large, deep convolutional neural network to classify the 1.2 million high-resolution images in the ImageNet ILSVRC-2010 contest into the 1000 different classes. On the test data, we achieved top-1 and top-5 error rates of 37.5% and 17.0%, respectively, which is considerably better than the previous state-of-the-art. The neural network, which has 60 million parameters and 650,000 neurons, consists of five convolutional layers, some of which are followed by max-pooling layers, and three fully connected layers with a final 1000-way softmax. To make training faster, we used non-saturating neurons and a very efficient GPU implementation of the convolution operation. To reduce overfitting in the fully connected layers we employed a recently developed regularization method called “dropout” that proved to be very effective. We also entered a variant of this model in the ILSVRC-2012 competition and achieved a winning top-5 test error rate of 15.3%, compared to 26.2% achieved by the second-best entry.

1. PROLOGUE

Four years ago, a paper by Yann LeCun and his collaborators was rejected by the leading computer vision conference on the grounds that it used neural networks and therefore provided no insight into how to design a vision system. At the

that were widely investigated in the 1980s. These networks used multiple layers of feature detectors that were all learned from the training data. Neuroscientists and psychologists had hypothesized that a hierarchy of such feature detectors would provide a robust way to recognize objects but they had no idea how such a hierarchy could be learned. There was great excitement in the 1980s because several different research groups discovered that multiple layers of feature detectors could be trained efficiently using a relatively straight-forward algorithm called backpropagation^{18, 22, 27, 33} to compute, for each image, how the classification performance of the whole network depended on the value of the weight on each connection.

Backpropagation worked well for a variety of tasks, but in the 1980s it did not live up to the very high expectations of its advocates. In particular, it proved to be very difficult to learn networks with many layers and these were precisely the networks that should have given the most impressive results. Many researchers concluded, incorrectly, that learning a deep neural network from random initial weights was just too difficult. Twenty years later, we know what went wrong: for deep neural networks to shine, they needed far more labeled data and hugely more computation.

2. INTRODUCTION

Current approaches to object recognition make essential use of machine learning methods. To improve their perfor-

Alex Krizhevsky, Ilya Sutskever, Geoffrey E. Hinton: *ImageNet classification with deep convolutional neural networks*. In: *Commun. ACM*. Band 60, Nr. 6, 24. Mai 2017

Zur Generativen KI

Nicht mehr nur klassifizieren, sondern erzeugen!

- Transformer Architektur
 - Geeignet für sequenzielle Information (Text, Audio, ...)
 - Für jedes Wort wird die Beziehung zu den anderen Wörtern ermittelt (Self Attention)
 - Die **Bank** steht am **Fluss**: Fluss ist wichtig, Geld kommt nicht vor, ist wohl eine Sitzbank
 - Ermöglicht „Verstehen“ von langen Texten und Beziehungen
- Grundlage für Text- und Bildgeneratoren (GPT)
 - Generative Pre-Trained Transformer
- In der ZfP nicht direkt relevant



ChatGPT: Generiere mir ein Symbolbild zum Thema KI mit blauer, eher dunklerer Farbstimmung

The background of the slide is a futuristic, digital-themed image. It features a central figure of a robot head, possibly a humanoid AI, with a glowing blue network of lines and nodes emanating from its head, suggesting a complex neural network or data processing. The overall color palette is dark blue and black, with bright blue highlights and glowing effects. The text is overlaid on this background.

ALLGEMEINE CHANCEN UND RISIKEN DER KI

Eine disruptive Entwicklung

AI

Gesellschaftliche und ethische Risiken

Wir könnten alle sterben!

- Desinformation und Vertrauensverlust
 - Deep Fakes, generierte Bilder, Troll-Bots
- Machtkonzentration
 - Erhebliche Rechenleistung erforderlich, nur wenige große Player können das leisten, Abhängigkeit entsteht
- Intransparente Entscheidungen
 - Kredite, Versicherungen, Bewerbungen, Strafverfolgung
- Arbeitsmarkt und Sinnverlust
 - Automatisierung betrifft plötzlich auch Knowledge-Worker und Kreative, Radiologen

Technische Risiken

Was kann schon schiefgehen!

- Intransparenz (Black Box)
 - Moderne Modelle sind „nicht erklärbar“, Milliarden Parameter
- Datenabhängigkeit und Bias
 - Falsche oder verzerrte Trainingsdaten ergeben falsche oder verzerrte Ergebnisse!
- Overfitting und schlechte Generalisierung
 - Modell versagt bei neuen, leicht veränderten Situationen
- Unzuverlässige Unsicherheitsabschätzung
 - KI weiß nicht, was sie nicht weiß (ein „inverser Sokrates!“)
- Emergentes Verhalten
 - Überraschende Fähigkeiten und Fehler!

Probleme mit der Intransparenz

Warum nochmal ist dieses Teil schlecht?

- Große KI Modelle sind intransparent
- „Explainable AI“ existiert nicht so richtig
 - Die KI kann zeigen, „wo“ sie hinschaut (Heatmaps, Feature Importance)
 - Aber nicht „warum“ sie da hingeschaut hat
 - Die Erklärungen sind oft vereinfachend und nachträglich
 - Zertifizierbarkeit ist praktisch nicht gegeben



Quelle: geeksforgeeks.org

Chancen der KI

Sie macht uns schneller.

- Medizin und Gesundheit
 - Bessere Diagnostik und Früherkennung, personalisierte Therapien
- Ein Produktivitätsverstärker
 - Abnehmen von „intellektueller Fleißarbeit“. Die Kontrolle bleibt vorerst beim Menschen
- Besserer Bildungszugang für viele (lange nicht alle!)
 - Individuelle Lernunterstützung
- Unterstützung bei globalen Herausforderungen
 - Bessere Klimamodelle

Künstliche Intelligenz bewährt sich am Universitätsklinikum Leipzig

Mittwoch, 21. Januar 2026



/picture alliance, Jan Woitas



Leipzig – Das Institut für Neuroradiologie am Universitätsklinikum Leipzig (UKL) bewertet den Einsatz von Künstlicher Intelligenz (KI) zur Unterstützung der Magnetresonanztomografie (MRT) positiv.

Sechs Monate nach der Implementierung eines zertifizierten KI-Tools zeigt sich demnach, dass die neue Hard- und Software nicht nur die räumliche Auflösung der gewonnenen Bilder verbessert, sondern auch die Untersuchung um bis zu 65 Prozent beschleunigen kann.

Das Problem der Verantwortung

Wer war schuld?

- Wenn die KI Entscheidungen trifft, wer verantwortet diese
 - Der Entwickler
 - Der, der die KI trainiert hat
 - Endanwender, der keine Ahnung hat wie sie funktioniert?
- Kann man das versichern?
- Ein vergleichbares Problem blockiert das autonome Fahren



UND WAS HAT DAS MIT ZFP ZU TUN?

Endlich kommt er zum Punkt!

AI

Neuronale Netze in der ZfP

AlexNet's Nachfahren

- Deep Learning
- Wunsch: Die KI sagt uns, welche Teile gut und welche schlecht sind, und am besten auch warum
- Erwartungshaltung oft:
 - Die KI kann das alleine herausfinden (kann sie nicht)
 - Die KI kann das besser als ein Mensch (kann sie manchmal)
 - Die KI kann das schneller (ist sie)

Herausforderungen

Probleme darf man ja nicht mehr sagen....

- CNNs trainieren mit wenigen oder schlechten Trainingsdaten
 - Nur ein paar Gut- und Schlechteile
 - Bias in den Daten, z.B. sehr viele Gutteile, sehr wenige Schlechteile mit nur einem Fehlertyp
 - Keine oder unzuverlässige Labels/Klassifizierungen
- Auswege, die ein wenig helfen
 - Vortrainierte CNN (Die ersten Ebenen sind recht universell)
 - Synthetische Trainingsdaten (nicht ganz ungefährlich)
 - Automatisches Labeln mit anderen Verfahren (Schwellen z.B.)
Aber dann kann das CNN auch nicht mehr als die etablierten Verfahren...
 - Aber: man müsste keine Schwellen etc. mehr einstellen

Was manchmal klappt...

It's the data, dude!

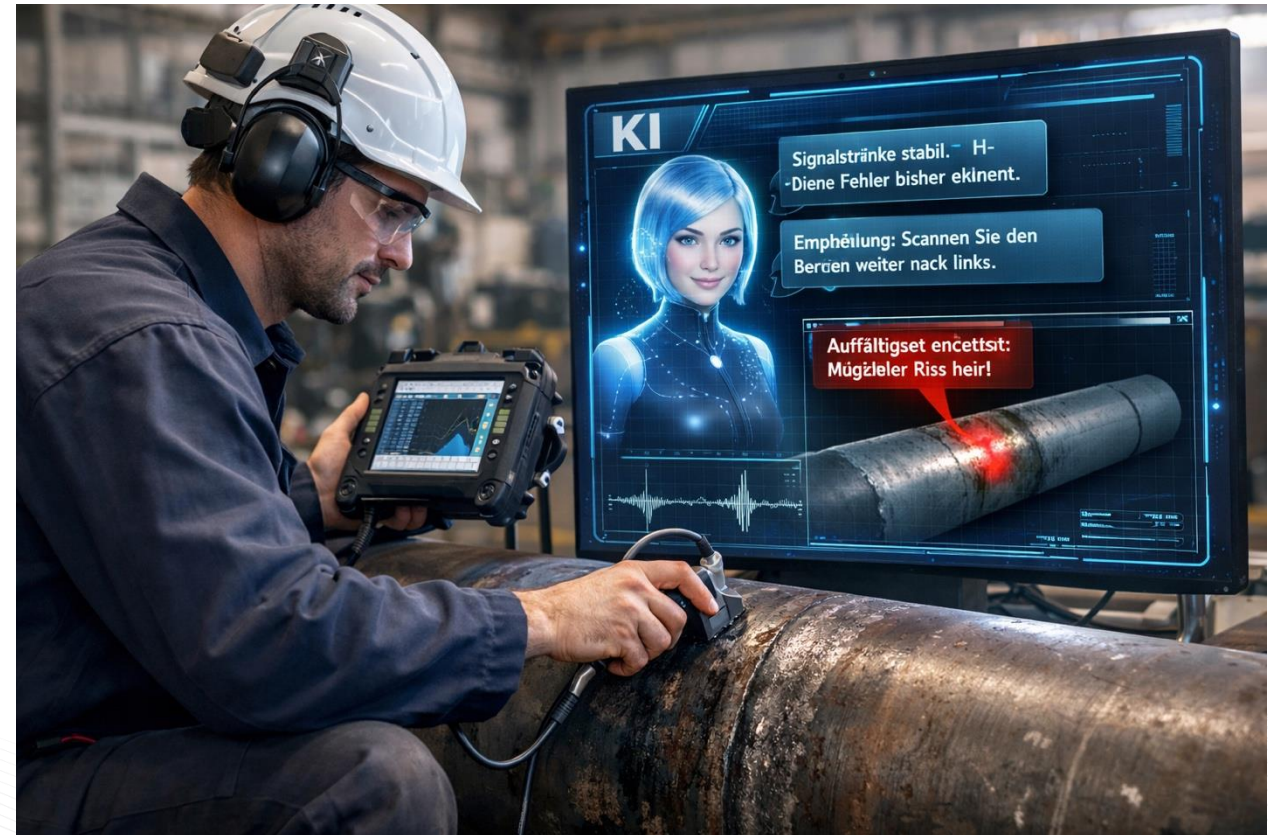
- Massenteile die sich in Geometrie und Material kaum unterscheiden
 - Spezifische Fehlerbilder
 - Mit vortrainierten Netzen kann man mit ein paar hundert Teilen zurechtkommen
 - Die können auch noch manuell gelabelt werden
 - Synthetische Trainingsdaten (nicht ganz ungefährlich)
 - Automatisches Taggen mit anderen Verfahren (Schwellen z.B.) Aber dann kann das CNN auch nicht mehr als die etablierten Verfahren...
- Und:
 - Wenn “Schlupf” zulässig ist (keiner stirbt)
 - Wenn keine Zertifizierung/POD Studie erforderlich ist
 - Wenn keine persönliche Verantwortlichkeit erforderlich ist



Aber zur Unterstützung des Prüfers ist es kein Problem

Oder?

- Oft gebrachtes Argument:
 - Die KI unterstützt den menschlichen Prüfer, um die Prüfsicherheit zu erhöhen
- Problem:
 - Wenn sie gut funktioniert, wird der automatisch immer weniger genau hinschauen. (Automation Bias)
 - Er ist dann „Schuld“, die KI (und ihr Hersteller) ist fein raus
 - Prüfer verlernen das Prüfen



Quelle: ChatGPT. Prompt: „Ein Bild, auf dem ein Mensch eine Zerstörungsfreie Prüfung durchführt und dabei von einer KI unterstützt wird“

Wo wir Deep Learning noch nicht einsetzen sollten

Und das kann sich schnell ändern....

- Fälle, wo bisher ein menschlicher Prüfer die Prüfentscheidung getroffen hat
- Bauteile, deren Versagen unmittelbar Menschenleben gefährden
 - Luftfahrt, Raumfahrt, Schiene (Radreifen), Sicherheitsbauteile am Auto
- Wenn wir genau verstehen müssen, wie die Prüfentscheidung zustande gekommen ist

Tim Reckmann, ccnull.de unter CC-BY 2.0 Lizenz, KI generiert

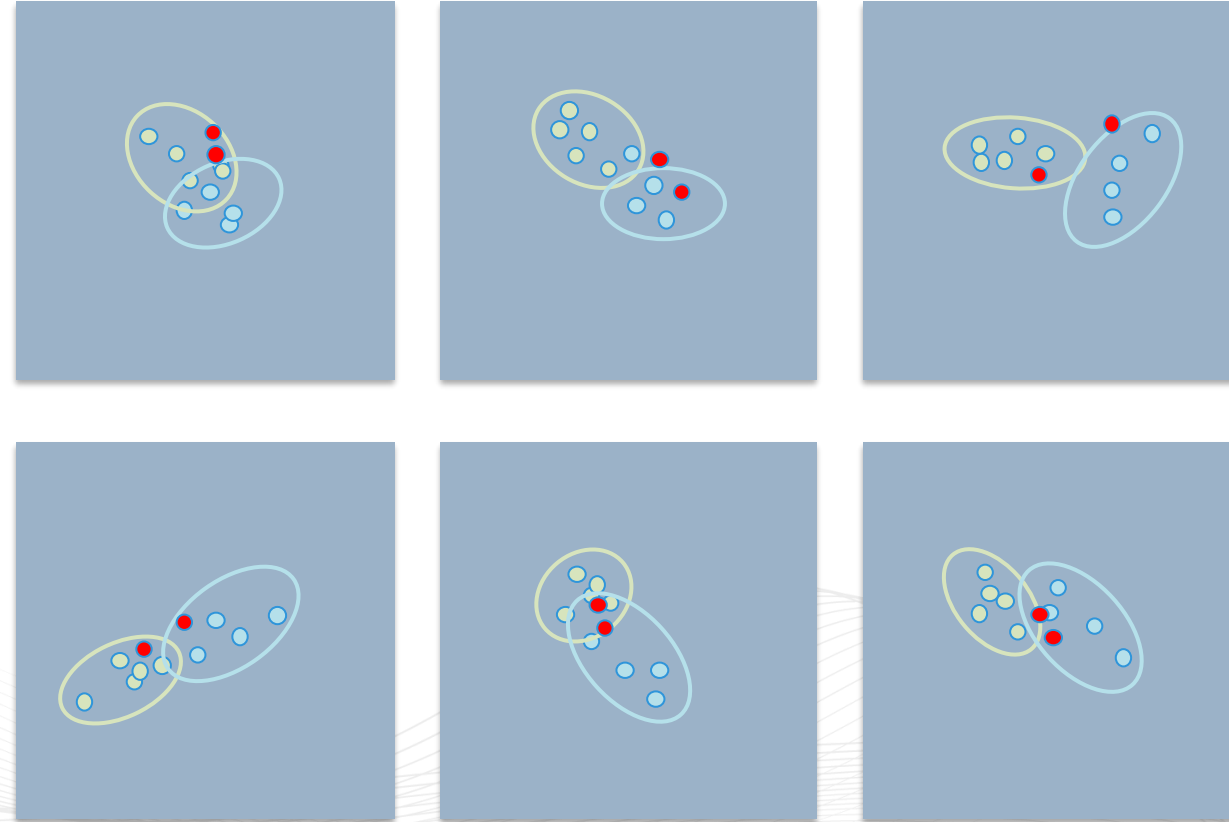


Alternativen! Aber: Ist das KI?

Oft können „ältere“ Machine Learning Verfahren eingesetzt werden

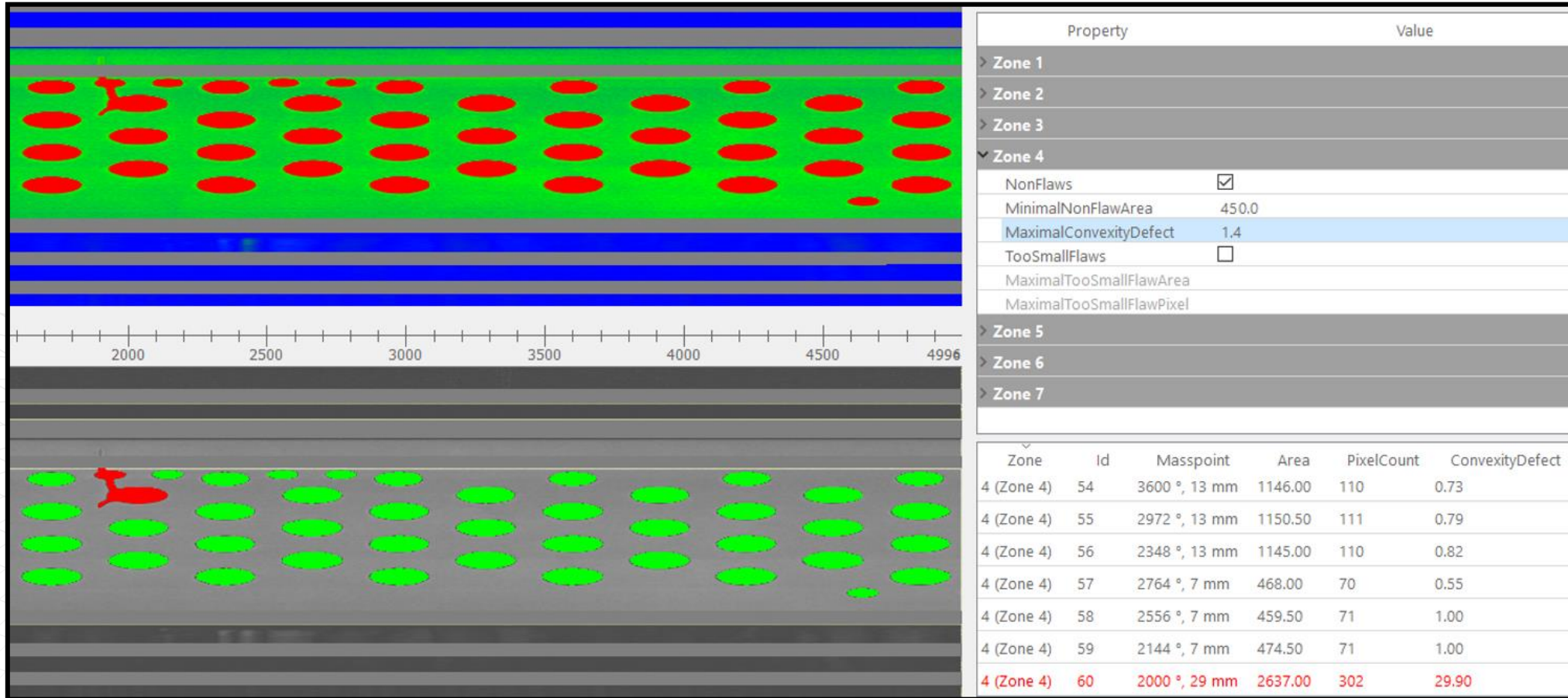
- Machine Learning: Das gibt's schon lange!
 - Sortierprüfung mit Gutteilen angelernt
 - „Automatische“ Signaloptimierung
 - Ausblenden von Störkonturen mit Feature Recognition
- Typische ML Verfahren
 - Lineare Regression
 - Support Vector Machines
 - k nearest neighbours
 - Entscheidungsbäume

Aber: Ist das „Intelligenz“?



„Klassische“ Bildverarbeitung

Auch ein bisschen Machine Learning

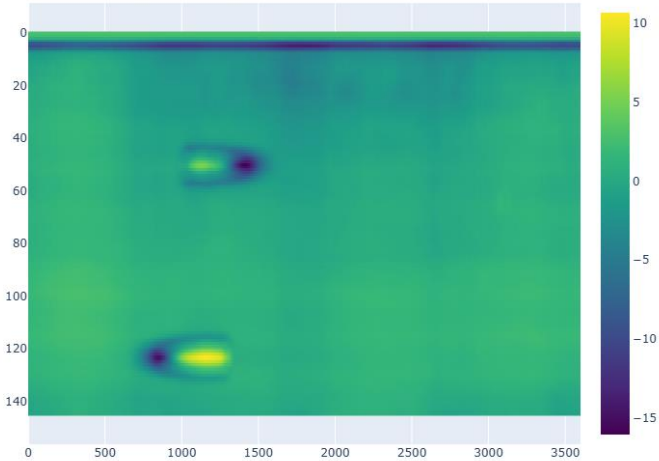


Beispiel für ML mit linearer Diskriminanzanalyse (LDA)

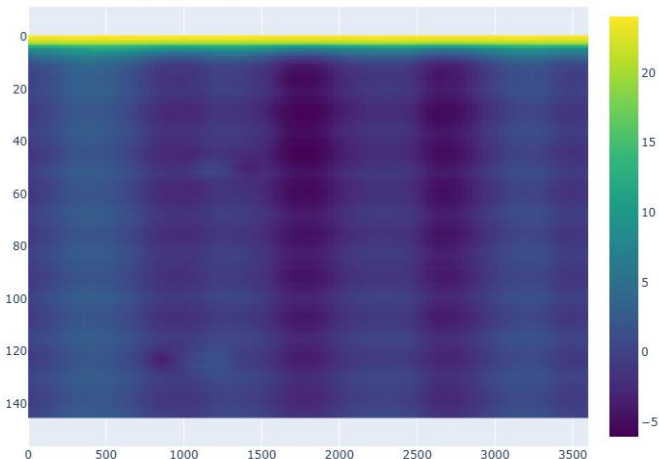
Beispiel: Schleifbranderkennung

Eingangsdaten

Feature: PCA2 (Box/Lasso)

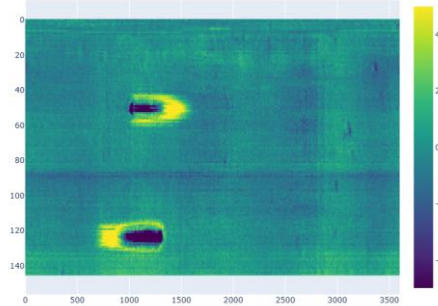


Feature: PCA1 (Box/Lasso)

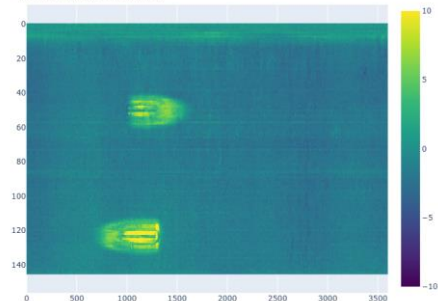


Sekundärmerkmale

Feature: LDA3 (Box/Lasso)

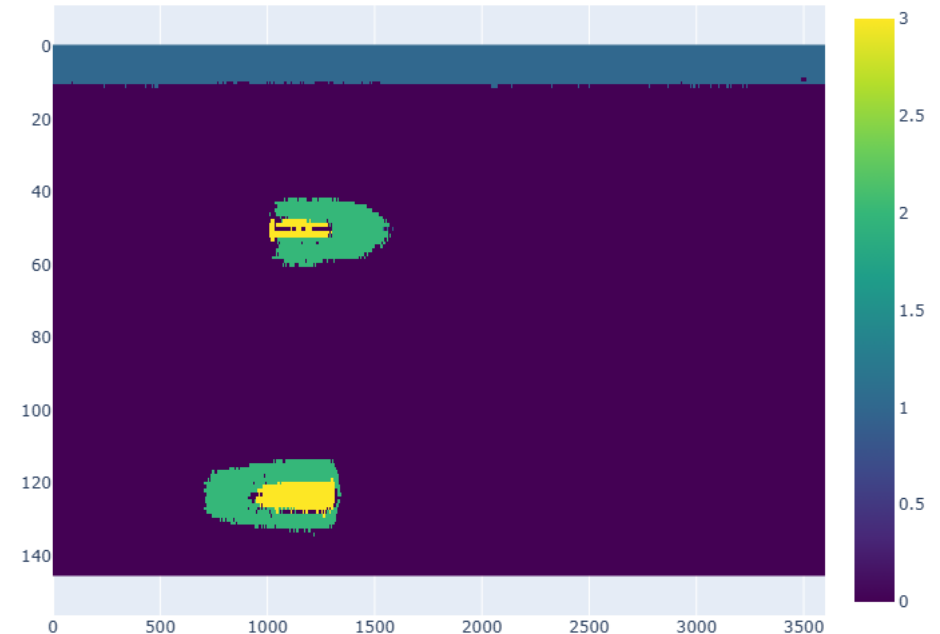


Feature: LDA2 (Box/Lasso)



Klassifizierung

Feature: LDA_pred (Box/Lasso)



Dr. Sargon Youssef, Rohmann GmbH

Fazit

Es gibt viel zu tun

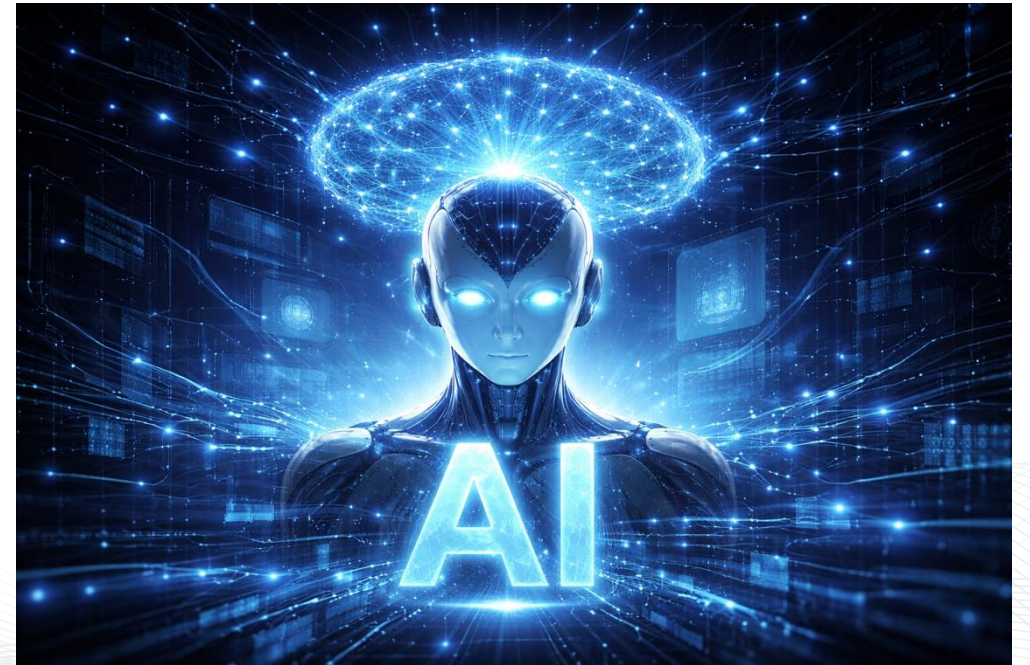
KI ist ein Werkzeug, kein Wundermittel.

Wir müssen lernen sie zu nutzen und ihre Grenzen kennen.

KI wird die Welt — also auch die Welt der ZfP — nachhaltig verändern

ROHMANN GMBH

Vielen Dank für Ihre Aufmerksamkeit!



Kontakt: gopalan@rohmann.de
sales@rohmann.de