

Auf den historischen Spuren der KI

Wie Maschinen Menschen austricksen

Bereiche, in denen KI unseren Alltag bereits erobert hat

Kapitel 1

Eine kurze Geschichte der KI

Um die Rolle der künstlichen Intelligenz (KI) in Wirtschaft und Gesellschaft vollständig zu verstehen, müssen wir erst einmal ihre faszinierende Geschichte nachvollziehen. Denn diese Spurensuche wirft nicht nur ein Licht auf die enormen Fortschritte der KI, sie verdeutlicht auch ihren Nutzen für das Marketing.

Die ersten Vorstellungen von künstlicher Intelligenz gehen auf die griechische Mythologie zurück, in der Talos, ein 2,44 Meter großer Riese aus Bronze, die Insel Kreta bewachte, um sie vor Piraten und anderen Eindringlingen zu schützen. Talos warf Felsbrocken auf Schiffe und patrouillierte jeden Tag auf der Insel. Der Legende nach wurde er schließlich besiegt, als ein Bronzenagel aus der Ferse seines Fußes entfernt wurde, sodass das *Ichor* (Blut der Götter) aus seinem Körper fließen konnte.

Von diesem Zeitpunkt an rankten sich zahlreiche Geschichten um automatisierte Wesen in der griechischen Mythologie, was die Fantasie von Wissenschaftlern, Mathematikern und Erfindern anregte. Die moderne Wissenschaft hat inzwischen einige dieser Mythen tatsächlich verwirklichen können. In diesem Kapitel stelle ich Ihnen diese Fortschritte vor, darunter den Turing-Test, maschinelles Lernen, Expertensysteme und generative KI.

Frühe technologische Fortschritte

Wissenschaftler führen die Anfänge der Automatisierung auf das 17. Jahrhundert und die Erfindung der *Pascaline* zurück, einer mechanischen Rechenmaschine. Dieses Gerät, das der französische Erfinder Blaise Pascal zwischen 1642 und 1644 konstruierte, verfügte über einen *kontrollierten Übertragsmechanismus*. Dieser erleichterte Rechenoperationen mit Addition und Subtraktion, indem er die Ziffer »1« effektiv in die nächste Spalte übertrug. Die Rechenmaschine arbeitete vor allem bei der Verarbeitung großer Zahlen effizient. Der deutsche Mathematiker Wilhelm Leibniz baute auf dieser Idee auf und erfand 1694 eine neue Rechenmaschine. Diese erweiterte das Konzept der Pascaline, indem sie alle vier

Grundrechenoperationen ermöglichte – Addition, Subtraktion, Multiplikation und Division (nicht nur Addition und Subtraktion). Damit boten diese Geräte erstmals einen Einblick in das Potenzial des mechanischen Denkens.

Wenn wir ins frühe 19. Jahrhundert vorspulen, stoßen wir auf das Jacquard-System, das von dem Franzosen Joseph-Marie Jacquard entwickelt wurde. Dabei wurden austauschbare Lochkarten verwendet, um das Weben von Stoffen und das Entwerfen komplizierter Muster zu bestimmen. Diese Lochkarten waren der Grundstein für spätere Entwicklungen in der Computertechnik. Mitte des 19. Jahrhunderts stellte der britische Erfinder Charles Babbage das erste Rechenggerät vor, die sogenannte *Analytical Engine*. Mithilfe von Lochkarten konnte diese Maschine eine Vielzahl von Berechnungen mit mehreren Variablen durchführen. Sie verfügte zudem über eine Reset-Funktion, wenn sie ihre Aufgabe erledigt hatte. Wichtig war auch, dass sie einen temporären Datenspeicher für komplexere Berechnungen hatte – eine entscheidende Funktion für jedes spätere System künstlicher Intelligenz (KI).

Ende der 1880er-Jahre erreichte die Entwicklung der KI mit der Entwicklung der Tabelliermaschine einen weiteren Meilenstein. Diese war vom amerikanischen Erfinder Herman Hollerith speziell zur Verarbeitung der Daten für die US-Volkszählung von 1890 konstruiert worden. Das elektromechanische Gerät verwendete Lochkarten zur Speicherung und Aggregation von Daten und verbesserte die Speicherkapazität der analytischen Maschine durch die Integration eines Akkumulators. Bemerkenswerterweise blieben modifizierte Versionen der Tabelliermaschine bis in die 1980er-Jahre betriebsbereit.

Alan Turing und die maschinelle Intelligenz

Viele Menschen betrachten Alan Turing, einen britischen Mathematiker, Logiker und Informatiker, als den Gründervater der theoretischen Informatik. Er ebnete den Weg für weitere Durchbrüche in der KI. Während des Zweiten Weltkriegs diente er in Bletchley Park, der Codeknacker-Einrichtung des Vereinigten Königreichs. Er spielte eine entscheidende Rolle bei der Decodierung von Nachrichten, die von der deutschen *Enigma-Maschine*, einem Gerät zur Codegenerierung, verschlüsselt wurden. Wissenschaftler und Historiker schreiben seiner Arbeit zu, dass sie den Krieg verkürzte und Millionen von Menschenleben rettete.

Turings wichtigste Innovation in Bletchley war die Entwicklung der *Bombe*, einer Maschine, die den Prozess zur Entschlüsselung von Nachrichten der Enigma-Maschine erheblich beschleunigte. Die Enigma verwendete eine Reihe rotierender Scheiben, um Klartextnachrichten in verschlüsselten Geheimtext umzuwandeln. Die Komplexität dieses Verschlüsselungsgeräts und der von ihm generierten verschlüsselten Nachrichten war auch darauf zurückzuführen, dass die Enigma-Benutzer die Einstellungen der Maschine täglich änderten. Für Großbritannien und alle Alliierten war es äußerst schwierig, den Code innerhalb des 24-Stunden-Fensters zu knacken – bevor die Einstellungen erneut geändert wurden. Die Turing-Bombe automatisierte den Prozess zum Identifizieren der Enigma-Einstellungen und analysierte verschiedene mögliche Kombinationen viel schneller als ein Mensch es jemals geschafft hätte. Diese Automatisierung ermöglichte es den Briten, die geheime deutsche Kommunikation regelmäßig zu entschlüsseln.



Obwohl die Einzelheiten dieses Codeknacker-Geräts viele Jahre lang unter Verschluss blieben, ist die Bombe eines der ersten Beispiele dafür, dass Technologie den Menschen bei Aufgaben, die traditionell menschliche Intelligenz erforderten, übertraf und sie effizienter und präziser ausführte.

Der Turing-Test im Jahr 1950

Kurz nach dem Zweiten Weltkrieg formulierte Turing in einem 1950 veröffentlichten Aufsatz mit dem Titel »Computing Machinery and Intelligence« die Idee, einen Standard zu definieren, ab dem man eine Maschine als intelligent bezeichnen könne. Er entwickelte das Experiment (heute *Turing-Test* genannt), das die Frage beantworten sollte: »Können Maschinen denken?« Die grundlegende Prämisse lautete dabei: Wenn ein Computer an einem Dialog mit einem Menschen teilnehmen kann und ein Beobachter nicht erkennen kann, welcher Teilnehmer Mensch und welcher ein Computer ist, dann kann man diesen Computer als intelligent bezeichnen.

Turings Test sah vor, dass ein menschlicher Prüfer Dialoge zwischen einem Menschen und einer Maschine beurteilen soll. Der Prüfer weiß zwar, dass einer der Teilnehmer eine Maschine ist, aber nicht, welcher. Um jegliche Verzerrung durch stimmliche Hinweise auszuschließen, schlug Turing vor, dass der Prüfer die Interaktionen auf ein reines Textmedium beschränkt. Wenn der Prüfer es schwierig fand, zwischen der Maschine und dem menschlichen Teilnehmer zu unterscheiden, bestand die Maschine den Test. Die Bewertung konzentrierte sich dabei nicht auf die Genauigkeit der Antworten der Maschine. Sondern darauf, wie wenig sich ihre Antworten von denen eines Menschen unterscheiden ließen.

Der Turing-Test: 1960er-Jahre und später

1966, lange nach Alan Turings Tod, entwickelte der deutsch-amerikanische Wissenschaftler Joseph Weizenbaum ELIZA, das erste Programm, das den Turing-Test zu bestehen schien. Viele Quellen bezweifeln zwar, dass es den Turing-Test bestehen konnte. Aber es war technisch in der Lage, einige Menschen davon zu überzeugen, sie würden mit menschlichen Operatoren sprechen. Das Programm funktionierte, indem es die getippten Kommentare eines Nutzers auf Schlüsselwörter hin untersuchte und daraufhin eine Regel ausführte, die die Kommentare veränderte. Das führte dazu, dass das Programm mit einem neuen Satz antwortete. Tatsächlich spiegelte ELIZA, wie viele Programme seitdem, ein Verständnis der Welt vor, ohne tatsächlich über reales Wissen zu verfügen.

Der amerikanische Psychiater Kenneth Colby ging 1972 noch einen Schritt weiter und entwickelte PARRY, das er als ELIZA mit Haltung beschrieb. Erfahrene Psychiater testeten PARRY in den frühen 1970er-Jahren mit einer Variante des Turing-Tests. Sie analysierten Texte von echten Patienten und von Computern, auf denen PARRY lief. Die Psychiater identifizierten die Patienten nur in 52 Prozent der Fälle richtig, eine Statistik, die mit zufälligen Vermutungen übereinstimmt.



Bis heute ist der Turing-Test eine prägnante, leicht verständliche Methode, um zu beurteilen, ob eine Technologie intelligent ist oder nicht. Indem der Test auf textbasierte Interaktionen beschränkt wurde und die Abfragen in natürlicher Sprache (Konversationsenglisch) erfolgen, konnte jeder leicht nachvollziehen, was der Test leistete, als Turing ihn erstmals vorstellte. Und indem er die Genauigkeit der Antwort nicht in den Vordergrund rückte, konzentrierte er den Test auf die Bewertung dessen, was Menschen wirklich menschlich macht.



Seit Alan Turing den Turing-Test erstmals präsentierte, haben sich Computer sprunghaft weiterentwickelt. Betrachten wir einmal die fortlaufenden Entwicklung intelligenter Technologie:

- ✓ **Noch im Jahr 2021** hatten Chatbots, die in weiten Teilen der Welt verfügbar waren, Mühe, den Turing-Test durchgängig zu bestehen. Dienste wie Siri von Apple, Alexa von Amazon und Googles Assistant konnten zwar in natürlicher Sprache mit uns sprechen, zogen sich aber bei einigen der einfachsten Fragen schnell auf Allgemeinplätze zurück. So könnte beispielsweise die Aufforderung »Beschreiben Sie sich selbst nur mit Farben und Formen!« die Antwort »Okay, ich habe im Internet Folgendes zum Beschreiben von Farben und Formen gefunden ...« hervorbringen.
- ✓ **Seit 2023** können die wichtigsten Chat-Schnittstellen von OpenAI, Google und anderen den Turing-Test bestehen. Diese schnelle Veränderung zeigen, wie schnell die technologischen Fortschritte im Bereich der KI sind und dass sich innerhalb von nur 24 Monaten dramatisch viel verändern kann.

Die Dartmouth-Konferenz von 1956

Die Dartmouth-Konferenz von 1956 gilt in der akademischen Gemeinschaft oft als Geburtsstunde der künstlichen Intelligenz (KI) als eigenständiges Forschungsgebiet. Die Konferenz fand im Sommer des Jahres am Dartmouth College in Hanover, New Hampshire, statt. Für einen Zeitraum von sechs bis acht Wochen fanden sich dort Koryphäen aus verschiedenen Disziplinen – Informatik, kognitive Psychologie, Mathematik und Ingenieurwissenschaften – unter einem Dach zusammen. Die von den Informatikern John McCarthy, Marvin Minsky und Nathaniel Rochester sowie dem Mathematiker Claude Shannon organisierte Konferenz zielte darauf ab, »jeden Aspekt des Lernens oder jedes andere Merkmal der Intelligenz« zu untersuchen, wie es in einem Schreiben der Konferenz hieß.

Die Dartmouth-Konferenz von 1956 stellt aus mehreren Gründen eine Zäsur dar. Sie war mehr als nur ein Sommertreffen von Intellektuellen; sie war ein bahnbrechendes Ereignis, das die Entwicklung der KI, wie wir sie heute kennen, prägte. Sie prägte den Begriff »KI« und war gleichzeitig die erste Community, die Forschungsrichtungen und damit Innovationen in diesem Bereich über Jahrzehnte hinweg vorangetrieben hat.

- ✓ **Prägen des Begriffs »künstliche Intelligenz« (KI):** Die Konferenz verlieh einem Forschungsbereich einen Namen, der bis dahin vage definiert war und sich interdisziplinär über Mathematik, Informatik, Ingenieurwissenschaften und verwandte Bereiche erstreckte. John McCarthy, einem der Organisatoren, wird die Einführung des Begriffs zugeschrieben. Das trug dazu bei, die zukünftige Richtung der Forschung zu bestimmen, weil sich sämtliche Wissenschaftler um einen zentralen Forschungsgegenstand versammeln konnten.
- ✓ **Katalysator für zukünftige Forschung:** Sie legte die Forschungsagenda für die kommenden Jahrzehnte fest. Während der Konferenz führten die Teilnehmer intensive Diskussionen, Brainstorming-Sitzungen und sogar Experimente im Frühstadium zu grundlegenden Themen im KI-Bereich durch. Die Teilnehmer wollten herausfinden, ob sie Maschinen so programmieren konnten, dass sie Aspekte der menschlichen Intelligenz simulierten.

Forschungsthemen waren beispielsweise:

- Problemlösung
- symbolisches Denken
- neuronale Netze
- Sprachverständnis
- lernende Maschinen

Sie entwickelten Programme zum Schachspielen, zum Beweis mathematischer Theoreme und zur Generierung einfacher Sätze.

- ✓ **Bereitstellen einer kollaborativen Plattform für interdisziplinäre Forschung:** Forscher, deren Wege sich sonst vielleicht nie gekreuzt hätten, führten nun bedeutungsvolle Dialoge und knüpften Beziehungen, die in den kommenden Jahren und Jahrzehnten zu bedeutenden Kooperationen führen würden. Diese Interdisziplinarität war entscheidend, um ein komplexes Problem wie die Simulation menschlicher Intelligenz angehen zu können. Denn dies erforderte Kenntnisse aus verschiedensten Bereichen wie Psychologie, Neurowissenschaften, Linguistik, Operations Research und Wirtschaftswissenschaften.
- ✓ **Wichtige Finanzierung und Aufmerksamkeit für den sich entwickelnden Bereich der KI:** Die Sichtbarkeit und Glaubwürdigkeit, die diese Veranstaltung mit sich brachte, führten zu erhöhten Investitionen in die KI-Forschung sowohl aus dem staatlichen als auch aus dem privaten Sektor. Diese finanzielle Unterstützung war für die Entwicklung von Laboren, akademischen Programmen und Forschungsprojekten von entscheidender Bedeutung.

Maschinelles Lernen und Expertensysteme entstehen

Nach der Dartmouth-Konferenz (siehe vorheriger Abschnitt) entstanden zwei wichtige Teilgebiete, die zu den Eckpfeilern der künstlichen Intelligenz wurden: maschinelles Lernen und Expertensysteme. Bei den *Expertensystemen* handelte es sich um regelbasierte Methoden, die auf vordefinierten, von Menschen erstellten Anweisungssätzen basierten. *Maschinelles Lernen* (zunächst als *selbstlernende Computer* bezeichnet) stellte einen radikalen Wandel dar. Die Herangehensweise zielte darauf ab, Systeme zu entwickeln, die aus Daten lernten, anstatt vorgegebenen Regeln zu folgen.

Maschinelles Lernen

Arthur Samuel, ein amerikanischer Pionier auf dem Gebiet der Computerspiele und der künstlichen Intelligenz, prägte 1959 offiziell den Begriff des *maschinellen Lernens*. Im Gegensatz zu traditionellen Computermethoden, die für jede Operation auf explizite Anweisungen angewiesen waren, konzentrierte sich das maschinelle Lernen auf die Entwicklung von Algorithmen, die in der Lage sind, aus vorhandenen Daten eigenständig Schlussfolgerungen abzuleiten. Diese Algorithmen verwenden statistische Techniken, um Muster zu erkennen, Entscheidungen zu treffen oder auf der Grundlage dieser Muster zukünftige Ergebnisse vorherzusagen.

In den 1960er-Jahren leistete die Raytheon Company einen bedeutenden Beitrag auf diesem Gebiet. Sie entwickelte ein frühes Lernmaschinensystem, das verschiedene Arten von Daten analysieren konnte, darunter Sonarsignale, Elektrokardiogramme und Sprachmuster. Die Maschine verwendete eine Form des *bestärkenden Lernens*, eine Untergruppe des maschinellen Lernens, bei der der Algorithmus weitere Aktionen durch Versuch und Irrtum ermittelt. Im Wesentlichen wurde das System für richtige Entscheidungen belohnt und für falsche bestraft. Menschen bedienten und optimierten das System, wobei sie einen Knopf drückten, um Fehler zu kennzeichnen und zu korrigieren. Diese Korrekturen ermöglichten es der Maschine, sich anzupassen und ihre Leistung im Laufe der Zeit zu verbessern.

Zu den wichtigsten herausragenden Merkmalen des maschinellen Lernens gehören:

- ✓ **Anpassungsfähigkeit:** Anstatt sich darauf zu verlassen, dass Menschen Problemlösungen manuell codieren, können Computer durch maschinelles Lernen eigene Lösungen finden, indem sie große Datensätze untersuchen. Diese Freiheit hat zu bahnbrechenden Anwendungen in verschiedenen Bereichen geführt. So sind Algorithmen des maschinellen Lernens beispielsweise die Basis großer Sprachmodelle und Computersystems, die es Computern ermöglichen, Objekte und Personen in Bildern und Videos zu identifizieren und zu verstehen.

Solche Systeme können

- menschenähnlichen Text generieren,
- mit unglaublicher Genauigkeit Tausende von Objekten erkennen und Spam-E-Mails filtern,
- menschliche Sprache in Echtzeit transkribieren und übersetzen.

Auf diese Themen gehe ich in den folgenden Kapiteln (beispielsweise in den Kapiteln 4 und 5) ausführlich ein.

- ✓ **Effiziente und skalierbare Lösungen:** Da die Entwicklung spezifischer Algorithmen für jede Erkennungs-, Filter- oder Generierungsaufgabe sowohl kostspielig als auch zeitaufwendig wäre, bietet maschinelles Lernen eine weitaus effizientere und *skalierbare Lösung* (was bedeutet, dass die Lösung Aufgaben aus riesigen Datensätzen ausführen kann, ohne dass die Kosten entsprechend steigen). Der datenbasierte Ansatz zur Lösungsfindung hat die Art und Weise revolutioniert, wie Technologen an Probleme herangehen und sie lösen. Und er hat komplexe Aufgaben automatisiert wie beispielsweise die Überprüfung von Social-Media-Inhalten auf Hassreden – eine Aufgabe, die Informatiker Computern vor einigen Jahren niemals zugetraut hätten.



Weil sich das maschinelle Lernen kontinuierlich weiterentwickelt, erwarten Experten, dass seine Auswirkungen und Relevanz in verschiedenen Bereichen weiter zunehmen werden. Beispiele für die Auswirkungen auf verschiedene Geschäftsbereiche finden Sie in Kapitel 2.

Untersuchung von Expertensystemen

In den späten 1960er-Jahren konzentrierten sich viele Forscher auf die Erfassung *domänen-spezifischen Wissens* und legten damit den Grundstein für *Expertensysteme*, also technische Systeme oder Computer, die die Rolle von Experten in einem bestimmten Bereich, wie etwa der Arzneimittelforschung, spielten. Diese Expertensysteme waren die Vorläufer der heutigen KI-Systeme. In den 1970er-Jahren entwickelten Forscher einige der ersten Expertensysteme, darunter DENDRAL (für die chemische Massenspektrometrie) und MYCIN (zur Diagnose bakterieller Infektionen). Diese Systeme erfassten das Wissen und die Denkfähigkeiten menschlicher Experten, um so unterschiedliche Ratschläge geben zu können, darunter einfache medizinische Diagnosen, aber auch Explorationsstrategien für den Mineralienabbau.

Die Expertensysteme funktionierten in engen Themenbereichen gut. Aber die Kosten und Schwierigkeiten bei der Pflege und Skalierung ihres regelbasierten Wissens schränkten ihre Nützlichkeit deutlich ein. Forschung und Entwicklung von Expertensystemen verliefen ungefähr so:

- ✓ **In den späten 1970er-Jahren** unterstützte das Tauwetter des KI-Winters (siehe folgenden Abschnitt) die breitere Einführung von Expertensystemen in verschiedenen Branchen, darunter im Gesundheitswesen, im Finanzwesen und in der Fertigung. Während dieser Zeit entwickelten Informatiker spezielle Tools zur Erweiterung ihrer Expertensysteme. Gleichzeitig wuchs der Nutzen dieser Systeme exponentiell.

- ✓ **In den 1990er-Jahren** wurden die Grenzen der Expertensysteme deutlich, insbesondere ihre Unfähigkeit, aus ihren Verarbeitungserfahrungen zu lernen oder ihre Leistung ohne eine externe Programmierung zu verbessern. Dieser Mangel führte zu einem Rückgang bei der Entwicklung eigenständiger Expertensysteme. Informatiker begannen, sie in größere, komplexere Computersysteme zu integrieren.
- ✓ **In jüngerer Zeit** erleben die Ideen, die den Expertensystemen zugrunde liegen, eine Art Revival, obwohl sie häufig in hybriden Formen auftreten, die maschinelles Lernen (siehe vorangegangenen Abschnitt) und andere datengesteuerte Techniken einbeziehen. Auch wenn nicht viele Unternehmen wegen ihrer Einschränkungen eigenständige Expertensysteme entwickeln und einsetzen, bleibt das Grundkonzept der Erfassung und Anwendung menschlicher Expertise in Computermodellen ein wesentlicher Bestandteil der KI. Umfassendere KI-Lösungen integrieren Expertensysteme sogar als Ergänzung zu anderen fortgeschrittenen Methoden (wie maschinelles Lernen und natürliche Sprachverarbeitung (Natural Language Processing, kurz NLP); mehr dazu finden Sie im Abschnitt »Weitere KI-Entwicklungen in den 1980er-Jahren« weiter hinten in diesem Kapitel).



Die Einführung von Expertensystemen war also ein wichtiger Moment in der Geschichte der künstlichen Intelligenz. Die Entwicklung von Expertensystemen war Vorreiter von Wissenstechniken, die Informatiker noch heute zum Training von KI-Systemen verwenden. Allerdings basieren die meisten KI-Tools inzwischen eher auf maschinellem Lernen als auf explizit programmierten Regeln, die menschliches Eingreifen erfordern, da es einfacher zu skalieren ist.

Ein KI-Winter bricht an

Nach dem Hype um künstliche Intelligenz in den 1960er- und frühen 1970er-Jahren wurden die Grenzen der frühen KI deutlich. Das führte zu einer geringeren Finanzierung und einem sinkenden Interesses, was als *KI-Winter* bezeichnet wird. Der Lighthill-Bericht, der für den British Science Research Council erstellt und 1973 veröffentlicht wurde, trug zu diesem KI-Winter bei. Der Bericht kritisierte den Mangel an praktischen Anwendungen und stellte das Potenzial der KI-Forschung infrage. Diese Kritik führte in mehreren Ländern zu einer Kürzung der staatlichen Förderung, darunter auch im Vereinigten Königreich.

Doch selbst während dieser Zeit der reduzierten Finanzierung wurde eine Forschung weiter verfolgt, die vor allem grundlegende technische Fähigkeiten wie Wahrscheinlichkeitschlussfolgerungen, neuronale Netzwerke und intelligente Agenten weiterentwickelte. Selbst in diesen wenig optimistischen Zeiten erzielten damit einige Informatiker wichtige Fortschritte, bis das maschinelle Lernen in den 1980er-Jahren mit ihren wegweisenden Innovationen eine neue Ära einläutete.



Die Lehren aus dem KI-Winter der 1970er-Jahre prägen bis heute die ethische Debatte um realistische versus übertriebene Behauptungen in der KI-Welt. Diese Debatte ist wichtiger denn je, da weltweit unterschiedliche Meinungen über die Versprechen und Gefahren der KI aufeinanderprallen.

Der Stanford Cart: Von den 60ern bis in die 80er-Jahre

Wenn man über die Geschichte der künstlichen Intelligenz (KI) spricht, kommt man um die Geschichte des Stanford Cart nicht herum. Dabei handelt es sich um einen ferngesteuerten vierrädrigen Wagen, der erstmals in den 1960er-Jahren vorgestellt wurde und später mit einer Kamera und einem Bordcomputer für Sicht und Steuerung ausgestattet wurde. Dieser Prototyp war einer der ersten Versuche, ein selbstfahrendes Fahrzeug zu bauen. Der Wagen, dessen Entwicklung über einen Zeitraum von 20 Jahren dauerte, diente als Plattform für die Forschung in den Bereichen Computervision, Wegplanung und autonome Navigation.

Die Entwicklung des Stanford-Cart-Projekts spiegelte nicht nur die Entwicklung von KI und Robotik innerhalb von 20 Jahren wider, sondern prägte auch die weitere Forschung. Das Projekt bleibt damit ein Beweis für die nachhaltige und iterative Entwicklung im Bereich der KI.

Die Entwicklungsstufen des Stanford Cart umfassen:

- ✓ **Fernsteuerung:** In den 1960er-Jahren war die erste Version des Wagens mit einer Fernsteuerungsfunktion ausgestattet. Bei der Entwicklung des Wagens damit zu beginnen, war äußerst sinnvoll: Denn der Wagen diente auch als Forschungsplattform, um das Problem der Fernsteuerung eines Mond-Rovers von der Erde aus zu untersuchen.
- ✓ **Selbstnavigation:** Anfang der 1970er-Jahre wurde der Wagen mit einer Kamera und einem Bordcomputer ausgestattet, sodass er einen Hindernisparcours bewältigen konnte. Er machte Fotos und berechnete dann auf der Grundlage dieser Bilder den besten Weg. Später in den 1970er-Jahren ermöglichten fortschrittlichere Computervision-Algorithmen dem Wagen, komplexe Umgebungen schneller zu erfassen, gleichzeitig wurden auch die Bildverarbeitungsfähigkeiten verbessert.
- ✓ **Komplexe Navigation in Echtzeit:** In den 1980er-Jahren konnte der Wagen Straßen folgen und Hindernissen in Echtzeit ausweichen. Das war größtenteils auf Verbesserungen bei der Hard- und Software zurückzuführen, verursacht durch eine stark gestiegene Computerleistung. Diese Fähigkeit gilt als ein wichtiger Meilenstein in der Entwicklung autonomer Fahrzeuge, die erst Jahrzehnte später in die kommerzielle Produktion gingen. Die gesteigerte Rechenleistung ermöglichte schnellere und komplexere Berechnungen, während fortschrittliche Algorithmen es dem Wagen ermöglichten, sekundenschnelle Entscheidungen zu treffen.



Als eine der ersten praktischen Anwendungen von KI in der Robotik demonstrierte der Stanford Cart, wie Computer mit der realen Welt interagieren können. Die Computerkomponenten, die visuelle Eingaben und Analysen ermöglichten, unterstrichen die potenziellen Vorteile einer ausgefeilten Bilderkennung und Szeneninterpretation. Die Robotik und autonomen Systeme zur Wegplanung und Hindernisvermeidung verwenden noch in heutiger Zeit algorithmische Techniken, die der Stanford Cart erstmals vorstellte.

Weitere KI-Entwicklungen in den 1980er-Jahren

Die 1980er-Jahre gelten als entscheidendes Jahrzehnt in der Entwicklung der künstlichen Intelligenz, geprägt durch enorme Fortschritte in verschiedenen Teilbereichen, insbesondere im maschinellen Lernen, in neuronalen Netzwerken und in der Verarbeitung von natürlicher Sprache. In dieser Zeit wurden grundlegende Fortschritte erzielt, die das Fundament für die KI-Technologien von heute legten.

Zu den wichtigsten Entwicklungen dieses Jahrzehnts gehören:

- ✓ **Backpropagation:** Die Einführung und Bekanntmachung des Backpropagation-Algorithmus zum Trainieren neuronaler Netzwerke. Vor der Backpropagation war das Trainieren komplexer neuronaler Netzwerke sehr rechenintensiv und wenig effektiv. Der Backpropagation-Algorithmus rationalisierte den Trainingsprozess, indem er den Fehler zwischen vorhergesagten und tatsächlichen Ergebnissen effizient berechnete und diesen Fehler dann über das Netzwerk zurückverteilte, um die interne Gewichtung anzupassen (die die Eingabedaten in den Hidden Layer des Netzwerks transformieren). Diese Innovation erleichterte das Trainieren mehrschichtiger neuronaler Netzwerke und ebnete den Weg für komplexere Architekturen und Anwendungen.
- ✓ **Deep Learning:** Ein Teilgebiet des maschinellen Lernens, das neuronale Netzwerke mit drei oder mehr Schichten verwendet. Forscher wie Geoffrey Hinton, Yann LeCun und Yoshua Bengio (an verschiedenen Universitäten tätig) waren in dieser Zeit von entscheidender Bedeutung, da sie den Grundstein für dieses Teilgebiet legten. Die mehrschichtigen neuronalen Netzwerke kamen in einer Reihe von Anwendungen zum Einsatz – von der Bild- und Spracherkennung bis zum Verständnis natürlicher Sprache. Das wiederum sollte später Innovationen bei der Automatisierung verschiedener Geschäftsprozesse vorantreiben.
- ✓ **Verarbeitung natürlicher Sprache (Natural Language Processing, kurz NLP):** Ursprünglich basierten die NLP-Systeme von Programmierern größtenteils auf von Hand erstellten Regeln. In den 1980er-Jahren kam es jedoch zu einer deutlichen Verlagerung hin zu statistischen Modellen. Dadurch wurden diese Systeme robuster und skalierbar. Gleichzeitig war damit der Weg frei für Ansätze, die auf maschinellem Lernen basierten und die NLP-Landschaft verändern sollten. Fortan waren komplexere Anwendungen wie Chatbots, Übersetzungsdienste und Tools zur Stimmungsanalyse möglich.
- ✓ **Robotik:** Das Jahrzehnt markierte auch den Beginn bedeutender Fortschritte in der Robotik, von denen viele auf den grundlegenden Konzepten der KI basierten. Das Stanford-Cart-Projekt beispielsweise (siehe vorheriger Abschnitt) war ein entscheidender Katalysator für die Forschung an autonomen Systemen.

Schnelle Fortschritte der KI in den 1980er-Jahren und darüber hinaus

Die bemerkenswerte Reise der künstlichen Intelligenz (KI) reicht von ihren mythologischen Inspirationen (Talos, der bronzene Riese in der griechischen Mythologie, der Kreta beschützte) bis hin zu bahnbrechenden Erfindungen wie Pascals Rechenmaschine (besprochen im Abschnitt »Frühe technologische Fortschritte« weiter vorn in diesem Kapitel) und Projekten wie dem Stanford Cart (siehe den Abschnitt »Der Stanford Cart: Von den 60ern bis in die 80er-Jahre« weiter vorn in diesem Kapitel). Die Fortschritte seit Anfang der 2010er-Jahre haben die KI-Landschaft dann grundlegend transformiert und die Art und Weise verändert, wie Menschen über die Rolle der Technologie in verschiedenen Bereichen denken – in Wirtschaft und Gesellschaft.

Ab den 1990er-Jahren führten rasante Fortschritte in verschiedenen Disziplinen der KI-Forschung zu erweiterten Fähigkeiten im Bereich des maschinellen Lernens und des Deep Learning. Andere Fortschritte verhalfen der KI zu der Fähigkeit, scheinbar intuitiv zu denken und menschenähnliche Inhalte zu generieren.

Maschinelles Lernen wird erwachsen

Zwischen den 1990er- und den frühen 2000er-Jahren entwickelte sich das maschinelle Lernen zur dominierenden Kraft in der KI-Entwicklung. (Eine Einführung in das maschinelle Lernen finden Sie im Abschnitt »Maschinelles Lernen« weiter vorn in diesem Kapitel.) In diesem KI-Bereich werden Algorithmen verwendet, um riesige Datensätze zu analysieren und damit Muster zu erkennen und Vorhersagen ohne explizit programmierte Regeln zu treffen. Getrieben durch eine deutliche Steigerung der Rechenleistung und der Datenverfügbarkeit führte das maschinelle Lernen zu neuen Anwendungsfällen im Bereich der *Computervision* (wo Computer Informationen aus Bildern, Videos und anderen Eingaben ableiten) und *Empfehlungssysteme* (Informationsfiltersysteme, die dem Nutzer die für ihn relevantesten Informationen vorschlagen).

Diese Fortschritte in der KI waren auch deshalb möglich, weil die KI-Engines Zugriff auf große Datensätze hatten. Die Modelle, die zur Analyse dieser Datensätze verwendet wurden, ahmten die Mustererkennung und Entscheidungsfindung menschenähnlich nach, indem sie statistische Beziehungen zwischen den Daten nutzten. Diese Entwicklungen zeigten, wie schnell ein KI-System selbstständig aus Daten lernen (extrapolieren) konnte, ohne dass ein Programmierer spezifische und explizite Anweisungen für dieses System codieren musste. Maschinelles Lernen ist bis heute das Herzstück der KI.

Ein entscheidendes Schachspiel

In den 1990er-Jahren kam es zu einem Wendepunkt in der Geschichte der KI, der die Fantasie der Menschen auf der ganzen Welt beflügelte. IBMs Deep Blue, ein Schachcomputer, besiegte 1997 den amtierenden Schachweltmeister Garri Kasparow. Obwohl Deep Blue damals noch nicht über ein modernes neuronales Netzwerk verfügte und sich stattdessen auf

Brute-Force-Heuristiken und spezielle Schachalgorithmen verließ, wurden grundlegende Techniken des maschinellen Lernens integriert, um Brettpositionen zu bewerten und das Spiel zu verbessern. Deep Blues Sieg war damit ein weiterer bedeutsamer Fortschritt für KI und maschinelles Lernen. Er

- ✓ hat bewiesen, dass eine Maschine einen Menschen bei einer Aufgabe, die komplexe Entscheidungen über viele Schritte hinweg erfordert, übertreffen kann.
- ✓ löste große Debatten über die Zukunft der KI und ihre möglichen Auswirkungen auf alle Facetten des Lebens aus. Diese Diskussionen wurden durch die Einführung der generativen künstlichen Intelligenz deutlich beschleunigt (siehe den Abschnitt »Inhalte erstellen mit generativer KI« weiter hinten in diesem Kapitel).
- ✓ unterstützte Kasparows Ansicht, dass Maschinen und Menschen gemeinsam viel mehr erreichen könnten als jeder von ihnen allein. Er führte den Begriff *Advanced Chess* für eine Schachform ein, bei der Menschen gemeinsam mit Computersystemen Schach spielen, und betonte, dass menschliche Intuition und maschinelle Berechnungen zusammen eine nahezu unschlagbare Kombination darstellen.



Kasparows Idee des fortgeschrittenen Schachs hatte einen nachhaltigen Einfluss darauf, wie wir heute über KI denken. Viele KI-Forscher betrachten fortgeschrittenes Schach als Vorläufer moderner Theorien über KI, die einem Menschen in verschiedenen Bereichen als Assistent dient. (Satya Nadella, CEO von Microsoft, hat diese Unterstützung allgemein als KI-Co-Pilot bezeichnet.) In den folgenden Kapiteln gehe ich auf die Rolle der KI als ergänzendes Werkzeug für Menschen im Bereich Wirtschaft und Marketing ein. In diesen Diskussionen können Sie die philosophischen Wurzeln dieses kooperativen Ansatzes, der sich aus Kasparows Erkenntnissen ableitet, deutlich erkennen.

Auf den Spuren der Deep-Learning-Revolution

In den letzten Jahren hat das Aufkommen des Deep Learning die Fähigkeiten und Genauigkeit von KI-Systemen deutlich verbessert. Deep Learning baut auf den Grundlagen des traditionellen maschinellen Lernens auf. Es verwendet neuronale Netzwerke mit mehreren Schichten – oft als *tiefe neuronale Netzwerke bezeichnet* –, um bei Aufgaben wie Bildklassifizierung, Spracherkennung und Verarbeitung natürlicher Sprache ein hohes Maß an Genauigkeit zu erreichen.

Was Deep Learning von früheren KI-Technologien unterscheidet, sind die Weiterentwicklung der Rechenleistung, die Verfügbarkeit riesiger Datensätze und die Verwendung komplexer Algorithmen, die neuronale Netzwerke mit mehreren Schichten optimieren. Diese mehrschichtige Architektur ermöglicht es dem System, komplexe Beziehungen in den Daten zu modellieren, was zu bemerkenswert präzisen Ergebnissen führt.

Deep-Learning-fähige Systeme

- ✓ **sind der Motor einer Vielzahl von KI-Anwendungen, die heute im Einsatz sind.** Deep Learning verändert die Automatisierung erheblich, indem es Systemen ermöglicht, komplexe analytische und prädiktive Aufgaben völlig autonom und ohne

menschliches Eingreifen auszuführen. Ob Sie digitale Sprachassistenten wie Siri oder Alexa, sprachgesteuerte TV-Fernbedienungen oder fortschrittliche Fahrerassistenzsysteme in modernen Autos verwenden: Deep Learning ist die Schlüsseltechnologie, die vielen dieser Innovationen zugrunde liegt.

- ✓ **versprechen, in den nächsten Generationen ihre übergreifende Intelligenz zu steigern.** Diese zukünftigen Systeme werden wahrscheinlich weniger Daten für ein effektives Lernen benötigen, auf immer ausgefeilteren Prozessoren arbeiten und immer fortschrittlichere Algorithmen verwenden. Menschen, die KI-Technologien entwickeln, verfolgen das Ziel, die künstliche Intelligenz näher an die Komplexität und Fähigkeiten des menschlichen Gehirns heranzuführen.



Obwohl es für Wissenschaftler und Programmierer noch Jahrzehnte dauern dürfte, bis sie *künstliche allgemeine Intelligenz* erreichen – einen Zustand, in dem die KI über Denk-, Lernfähigkeiten und einen Verstand verfügt, die denen des Menschen ähneln –, ist Deep Learning zweifellos ein wichtiger Schritt auf diesem Weg.

Intuition beweisen im Zeitalter der KI

Der Turing-Test warf die grundlegende Frage auf: »Können Maschinen denken?« Man begann darüber nachzudenken, ob Menschen bei einer textbasierten Interaktion zwischen einer Maschine und einem Menschen unterscheiden können. (Informationen zum Turing-Test finden Sie im Abschnitt »Alan Turing und maschinelle Intelligenz« weiter vorn in diesem Kapitel.) Diese Frage schien 2016 durch den bahnbrechenden Sieg von AlphaGo über Lee Sedol in einer Partie Go eine endgültige Antwort zu finden.

AlphaGo war die Idee von DeepMind, einem britischen KI-Unternehmen, das später von Google aufgekauft wurde. Anders als herkömmliche KI-Programme wurde AlphaGo speziell dafür entwickelt, das Spiel Go zu meistern, ein altes Brettspiel, dessen Komplexität die von Schach bei Weitem übertrifft. Obwohl das Spiel einfache Regeln hat, verleiht ihm die schiere Anzahl der möglichen Züge eine astronomische Komplexität. Top-Go-Spieler wie Lee Sedol, eine führende Persönlichkeit in der Welt des Go, werden für ihre Intuition, Kreativität und analytischen Fähigkeiten verehrt.

Zur Vorbereitung auf das Duell mit Lee Sedol im Jahr 2016 wurde AlphaGo einem rigorosen Training unterzogen. Dabei kam eine Kombination aus Methoden des maschinellen Lernens, darunter Deep Learning, und anderen Algorithmen wie der wahrscheinlichkeitsbasierten Monte Carlo Tree Search zum Einsatz. Das Programm analysierte Tausende von historischen Go-Spielen und, was vielleicht noch beeindruckender ist, verfeinerte seine Fähigkeiten, indem es zahllose Spiele gegen sich selbst spielte. Dieses Selbstspiel ermöglichte es AlphaGo, verschiedene Strategien und Taktiken zu simulieren und so seine eigenen Spielfähigkeiten zu verbessern.

Als AlphaGo Lee Sedol in einer Serie von fünf Spielen besiegte, wurde die globale KI-Community hellhörig und nahm zwei überraschende Erkenntnisse wahr:

- ✓ **Der unerwartete Einfallsreichtums der KI:** Vor allem überraschte AlphaGos Fähigkeit, scheinbar kreative und intuitive strategische Entscheidungen zu treffen – Eigenschaften,

von denen viele annahmen, dass sie ausschließlich menschlicher Wahrnehmung vorbehalten seien. Sergey Brin von Google – dessen Unternehmen DeepMind übernommen hat – war beim dritten Spiel in Seoul dabei und sagte im Anschluss: »Wenn man wirklich großartigen Go-Spielern beim Spielen zusieht, hat das auch etwas von Schönheit. Ich bin begeistert, dass es uns gelungen ist, diese Art von Schönheit in unsere Computer zu bringen.«

- ✓ **Die tiefgreifenden Fähigkeiten und das Zukunftspotenzial der KI:** Der Sieg von AlphaGo war mehr als nur ein technologischer Meilenstein; er führte zu einem Paradigmenwechsel, der das Bewusstsein für KI bei Führungspersönlichkeiten aus den verschiedensten Sektoren schärfte – von Wissenschaftlern und Politikern bis hin zu Wirtschaftsführern und der breiten Öffentlichkeit.



Dieses historische Ereignis, bei dem AlphaGo einen nahezu unschlagbaren menschlichen Go-Spieler besiegte, war ein unwiderlegbarer Beweis für die Fortschritte im Bereich Deep Learning. Es zeigte, dass künstliche Intelligenz tatsächlich Aufgaben ausführen kann, von denen viele zuvor dachten, sie seien nur mit menschlicher Intelligenz ausführbar.

Inhalte erstellen mit generativer KI

Die Fortschritte im Bereich der KI nach 2010 führten zu bahnbrechenden Innovationen, insbesondere bei der Entwicklung *generativer Modelle* (die neue synthetische Daten wie Texte oder Bilder generieren können). In den 2020er-Jahren fanden generative Modelle Anwendung in einer Vielzahl von Bereichen – von Kunst und Unterhaltung bis hin zu wissenschaftlicher Forschung und der Entdeckung von Medikamenten. Zwei spezifische Entwicklungen bildeten die notwendige Grundlage für die Weiterentwicklung generativer Modelle:

- ✓ **Generative Adversarial Networks (GANs):** GANs wurden 2014 von dem Computerforscher Ian Goodfellow und seinen Kollegen vorgestellt und waren in der Lage, unglaublich realistische Bilder, Texte und andere Datentypen zu generieren. Dieser bedeutende Fortschritt bot ein robustes Framework für die Generierung komplexer, hochwertiger digitaler Assets. Spätere Weiterentwicklungen bei GANs führten zu Modellen wie zum Beispiel StyleGAN, das hochauflösende, äußerst realistische Bilder generieren kann.
- ✓ **Die Transformer-Architektur:** Ursprünglich für die Verarbeitung natürlicher Sprache entwickelt, wurde sie später für generative Aufgaben angepasst. Diese Anpassung an generative Aufgaben gipfelte in Modellen wie der GPT-Reihe von OpenAI, die menschenähnlichen Text generieren kann.

Ich befasse mich ausführlich mit generativen KI-Modellen in Kapitel 8.