

Was war die erste künstliche Intelligenz wirklich?

Category: KI & Automatisierung

geschrieben von Tobias Hager | 12. Dezember 2025



Was war die erste künstliche Intelligenz wirklich?

Die Frage mag einfach erscheinen, doch die Antwort ist alles andere als trivial. Während wir heute von Chatbots, Deep Learning und autonomen Systemen sprechen, hat die Geschichte der künstlichen Intelligenz (KI) ihre Wurzeln in den dunkelsten Ecken der Computerentwicklung. Doch was genau war eigentlich die erste echte KI? War es ein Algorithmus, ein Programm oder eine Maschine? Und vor allem: Wann hat KI wirklich angefangen, zu existieren – oder war das nur ein Mythos? Bereit, die schmutzige Wahrheit hinter den glanzvollen Marketing-Claims zu entdecken? Dann schnall dich an, denn wir tauchen tief in die technische Geschichte der ersten künstlichen Intelligenz ein – und lassen keine Frage offen.

- Definition: Was bedeutet eigentlich „künstliche Intelligenz“ im historischen Kontext?
- Die Pioniere: Von Alan Turing bis John McCarthy – wer hat was wirklich gemacht?
- Der erste KI-Algorithmus: Was war er, und warum war er so bedeutend?
- Technische Meilensteine: Von ELIZA bis zum ersten Expertensystem
- Mythen und Missverständnisse: Wann war wirklich die erste KI?
- Warum die erste KI noch kein echtes Bewusstsein hatte – und was das bedeutet
- Der technische Aufbau der allerersten KI: Hardware, Software und Methodik
- Was wir heute aus der frühen KI-Geschichte lernen können
- Der Einfluss auf moderne KI-Entwicklung: Kontinuität oder Bruch?
- Fazit: War die erste KI nur ein erster Schritt – oder ein falscher?

In einer Welt, in der KI fast schon zum Alltag gehört, ist es schwer, sich vorzustellen, dass alles einmal mit einem simplen Programm begann. Aber genau hier liegt der Clou: Die erste künstliche Intelligenz war kein Hochleistungsalgorithmus, kein neuronales Netz oder maschinelles Lernen im heutigen Sinn. Es war vielmehr eine Idee, ein Konzept, das sich erst über Jahrzehnte zu dem entwickelt hat, was wir heute als KI kennen. Die Geschichte ist voll von Mythen und Halbwahrheiten, aber nur eines ist sicher: Der Weg war lang, steinig und voll von technischer Ignoranz, aber auch genialer Durchbrüche.

Was bedeutet eigentlich

„künstliche Intelligenz“ im historischen Kontext?

Der Begriff „künstliche Intelligenz“ wurde erstmals 1956 auf der Dartmouth Conference geprägt – eine Art Gründungsmythos für die moderne KI-Forschung. Doch was genau versteckte sich damals hinter dieser Bezeichnung? Es ging vor allem um die Vorstellung, dass Maschinen menschliches Denken nachahmen und sogar übertreffen könnten – zumindest in bestimmten Bereichen. Die Definition war damals noch vage, aber sie war ambitioniert: Ein Computer sollte in der Lage sein, Probleme zu lösen, zu lernen und Entscheidungen zu treffen – alles auf eine Art und Weise, die an menschliche Intelligenz erinnert.

Viel wichtiger ist, dass die frühe KI-Definition stark von symbolischer KI geprägt war. Das bedeutete, dass man versuchte, menschliches Wissen in Form von Regeln, Fakten und logischen Schlüssen in Software zu gießen. Das Ziel war, Maschinen zu bauen, die mit einer Art Wissensbasis arbeiten und auf dieser Basis Schlussfolgerungen ziehen konnten. Damit war die Grundlage für die erste echte KI gelegt: eine regelbasierte, symbolische KI, die auf logischen Prinzipien fußte und keine neuronalen Netze oder Deep Learning kannte.

Diese frühe Definition hat bis heute Bestand – nur die technischen Methoden haben sich radikal verändert. Doch im Kern ist KI immer noch die Fähigkeit einer Maschine, Aufgaben zu bewältigen, die menschliche Intelligenz erfordern. Die Frage ist nur: Wann wurde diese Fähigkeit zum ersten Mal umgesetzt?

Die Pioniere: Von Alan Turing bis John McCarthy – wer hat was wirklich gemacht?

Ohne Alan Turing, den britischen Mathematiker und Logiker, wäre die Geschichte der KI kaum denkbar. 1950 veröffentlichte er sein legendäres Papier „Computing Machinery and Intelligence“, in dem er die berühmte Turing-Test-Idee vorstellte. Dieser Test sollte feststellen, ob eine Maschine in der Lage ist, menschliches Verhalten so gut zu imitieren, dass ein menschlicher Gesprächspartner keinen Unterschied erkennen kann. Turing war zwar kein KI-Entwickler im engeren Sinn, aber seine theoretischen Überlegungen legten den Grundstein für die spätere Praxis.

Doch die eigentliche Geburtsstunde der künstlichen Intelligenz als Forschungsdisziplin schlug 1956, als John McCarthy, Marvin Minsky, Nathaniel Rochester und Claude Shannon die Konferenz an der Dartmouth College organisierten. McCarthy, der Begründer des Begriffs „Artificial Intelligence“, sah darin die Chance, Maschinen zu bauen, die durch

symbolische Regeln intelligentes Verhalten zeigen. Seine Programmiersprache Lisp wurde zum Standard für frühe KI-Experimente und ersetzte die experimentellen Sprachen der Zeit.

In den folgenden Jahren entstanden die ersten symbolischen Expertensysteme, die versuchten, menschliches Fachwissen in Software zu kodieren. Diese Systeme waren regelbasiert, erklärbar und konnten auf vordefinierte Fragen antworten – doch sie waren auch spröde und schwer skalierbar. Trotzdem markierten sie einen Meilenstein in der technischen Entwicklung der künstlichen Intelligenz.

Der erste KI-Algorithmus: Was war er, und warum war er so bedeutend?

Der erste echte KI-Algorithmus, der breite Aufmerksamkeit erlangte, war das sogenannte General Problem Solver (GPS), entwickelt von Herbert Simon und Allen Newell. GPS war ein Versuch, eine universelle Problemlösungsmaschine zu bauen, die in der Lage war, jede Art von Problem zu lösen, für die eine formale Repräsentation existierte. Angetrieben wurde GPS von der sogenannten „Means-End-Analysis“, einem heuristischen Algorithmus, der darauf abzielte, den Problemraum schrittweise zu reduzieren.

GPS war bahnbrechend, weil es bewies, dass Maschinen durch formale Regeln und Heuristiken in der Lage sind, komplexe Probleme zu lösen. Es war kein reines Rechenmodell, sondern eine Art Vorläufer der modernen KI-Planungsalgorithmen. Allerdings war GPS extrem limitiert, da es nur auf klar definierte Probleme angewandt werden konnte und keine Lernfähigkeit besaß.

Dennoch markierte GPS den Beginn einer Ära, in der Künstliche Intelligenz nicht nur reine Datenverarbeitung war, sondern auf Problemlösungsstrategien basierte. Es war der erste Algorithmus, der zeigte, dass Maschinen mit den richtigen Strategien menschliche Denkprozesse nachahmen können – zumindest in eingeschränktem Rahmen.

Technische Meilensteine: Von ELIZA bis zum ersten Expertensystem

In den 1960er Jahren entstand ELIZA, das erste Programm, das tatsächlich als „Chatbot“ bezeichnet werden kann. Entwickelt von Joseph Weizenbaum, simulierte ELIZA einen Rogerian-Therapeuten, indem sie einfache Mustererkennung und Text-Substitution nutzte. Obwohl ELIZA kein echtes Verständnis hatte, war sie der erste Schritt in Richtung Interaktion zwischen

Mensch und Maschine. Sie zeigte, dass einfache regelbasierte Systeme in der Lage sind, auf menschliche Eingaben zu reagieren – eine Art frühes Vorläufer-Interface für KI.

Gleichzeitig entstanden die ersten Expertensysteme wie DENDRAL und MYCIN. Diese Systeme kodierten spezialisiertes Wissen in Form von Regeln und konnten Diagnosen stellen oder wissenschaftliche Hypothesen generieren. Sie waren die ersten funktionierenden KI-Anwendungen, die in realen Szenarien genutzt wurden – allerdings nur in engen Fachgebieten. Hier zeigte sich schon die Schwäche: Die Systeme waren zwar intelligent in ihrer Domäne, aber kaum flexibel.

Technisch gesehen basierten diese Systeme auf einer Kombination aus Wissensdatenbanken, Inferenzmaschinen und regelbasierten Engines. Für die damalige Zeit waren sie ein Quantensprung, denn sie bewiesen, dass Maschinen komplexe Aufgaben lösen können, die menschliches Fachwissen erfordern. Doch der technische Anspruch war hoch, und die Systeme waren wartungsintensiv und schwer skalierbar.

Mythen und Missverständnisse: Wann war wirklich die erste KI?

Viele glauben, die erste echte KI sei ELIZA gewesen. Falsch. ELIZA war lediglich eine regelbasierte Textmaschine ohne echtes Verständnis. Andere nennen den „Logic Theorist“ von Herbert Simon und Allen Newell – der erste Versuch, Logik und Problemlösung in Software zu packen. Doch auch dieser war noch weit von einem autonomen, lernfähigen System entfernt.

Der Mythos, die erste KI sei ein autonomes, lernendes System wie heute, ist schlichtweg falsch. Die frühe KI war alles andere als intelligent im heutigen Sinne. Sie war eine Ansammlung von Regeln, Heuristiken und Logik – oft schwer wartbar und auf eng umrissene Aufgaben beschränkt. Der technische Fortschritt war eher eine Aneinanderreihung von Speziallösungen, die nach und nach zusammenkamen.

Nur die technologische Entwicklung der letzten Jahre – neuronale Netze, Deep Learning, Reinforcement Learning – hat das Bild verändert. Doch die Grundidee, dass KI mehr ist als nur ein Programm, ist schon viel älter. Es war immer eine Vision, ein Ziel, das nur langsam Realität wurde.

Der technische Aufbau der

allerersten KI: Hardware, Software und Methodik

Die allererste KI lief auf den klassischen Großrechnern der 1950er und 1960er Jahre. Diese Systeme waren massiv, teuer und extrem begrenzt in ihrer Rechenleistung. Die Software basierte auf symbolischer Logik, Regelwerken und heuristischen Algorithmen. Das Betriebssystem war meist eine frühe Version von Fortran oder Lisp, die speziell für KI-Experimente angepasst wurden.

Der technische Kern war eine Kombination aus Inferenzmaschinen, Wissensbasen und Suchalgorithmen. Die Inferenzmaschine führte logische Schlussfolgerungen aus den Regeln und Fakten aus. Die Wissensbasen waren meist handgeschriebene Datenbanken, die Expertenwissen kodierten. Der Problemraum wurde durch Suchverfahren durchforstet, um Lösungen zu finden. Eine typische technische Herausforderung war die begrenzte Rechenleistung, was die Problemlösungsfähigkeit einschränkte.

Die Hardware war noch weit von den heutigen Standards entfernt – keine GPUs, kein Cloud-Computing, keine Massenspeicher. Alles musste auf einzelnen Großrechnern laufen, die nur die wichtigsten Funktionen unterstützten. Trotzdem legten diese Systeme den Grundstein für alles, was folgte. Ohne die technische Basis der ersten KI-Programme gäbe es heute kein Deep Learning oder neuronale Netze.

Was wir heute aus der frühen KI-Geschichte lernen können

Die wichtigste Lektion ist, dass KI kein plötzlicher Durchbruch, sondern ein langer, schmerzhafter Entwicklungsprozess ist. Die frühen Systeme waren zwar technisch limitiert, aber sie bewiesen, dass Maschinen in der Lage sind, bestimmte Aufgaben menschenähnlich zu lösen – wenn auch nur eingeschränkt. Sie haben gezeigt, dass formale Logik, Symbolik und Wissensrepräsentation zentrale Bausteine sind, die heute noch ihre Bedeutung haben.

Gleichzeitig offenbaren die frühen KI-Experimente die Grenzen: Rein regelbasierte Systeme sind schwer wartbar, wenig flexibel und kaum skalierbar. Die Erkenntnis, dass maschinelles Lernen und neuronale Netze eine völlig neue Dimension eröffnen, kam erst später. Doch die Grundidee – Maschinen, die denken – ist schon seit den Anfängen präsent.

Dieses historische Verständnis hilft uns heute, kritischer mit den aktuellen Hype-Begriffen umzugehen. Nicht jede „KI“-Lösung ist wirklich intelligent. Oft sind es nur raffinierte Mustererkennung oder statistische Modelle. Die Grundlagen der symbolischen KI sind jedoch nach wie vor relevant und bilden die Basis für Hybridansätze, die heute wieder auf dem Vormarsch sind.

Der Einfluss auf moderne KI-Entwicklung: Kontinuität oder Bruch?

Die heutige KI-Landschaft ist geprägt von Deep Learning, neuronalen Netzen und massivem Datenwachstum. Doch der Einfluss der frühen symbolischen KI ist unübersehbar. Viele moderne Ansätze sind Hybridmodelle, die Regeln, Logik und Lernen kombinieren. Das zeigt, dass die Geschichte der KI kein linearer Fortschritt ist, sondern ein komplexes Zusammenspiel verschiedener Paradigmen.

Der Bruch kam erst mit der Verfügbarkeit großer Datenmengen und leistungsfähiger Hardware. Damit wurde es möglich, Muster in Daten zu erkennen, die vorher unvorstellbar waren. Doch die Grundidee, Maschinen mit menschlicher Intelligenz auszustatten, ist älter als die meisten denken. Es ist eine iterative Entwicklung, bei der die alten Konzepte wiederentdeckt und erweitert wurden.

Für die Zukunft heißt das: Die erste KI war nur der Anfang. Heute bauen wir auf den Grundlagen der symbolischen KI auf, erweitern sie durch statistische Methoden und integrieren adaptive Lernprozesse. Doch ohne das Verständnis ihrer Ursprünge ist kein Fortschritt wirklich verständlich.

Fazit: War die erste KI nur ein erster Schritt – oder ein falscher?

Die Antwort auf die Frage, was die erste künstliche Intelligenz wirklich war, hängt stark von der Definition ab. War es ELIZA? War es GPS? Oder die symbolische KI? Fakt ist: Die erste echte KI war keine smarte Maschine, sondern ein Konzept, eine Vision. Sie war eine Ansammlung von Regeln, Logik und Methodik – weit entfernt von den autonomen, lernenden Systemen, die heute Schlagzeilen machen.

Doch gerade diese frühen Versuche haben den Grundstein gelegt für alles, was noch kommen sollte. Sie zeigten, dass Maschinen Denkprozesse nachahmen können – wenn auch nur in engen Grenzen. Und sie offenbarten die technischen und theoretischen Herausforderungen, die bis heute unser Verständnis prägen. War die erste KI also nur ein erster Schritt? Absolut. War sie ein Irrweg? Keineswegs. Sie war der Ausgangspunkt für eine Revolution, die noch lange nicht vorbei ist.

Wer heute in der KI-Forschung unterwegs ist, sollte die Geschichte kennen. Denn nur wer die Ursprünge versteht, kann die Zukunft sinnvoll gestalten. Und

die ist – wie so oft – eine Mischung aus alten Konzepten und neuen Technologien. Die erste künstliche Intelligenz war nur der Anfang – der Anfang einer Reise, die noch längst nicht zu Ende ist.