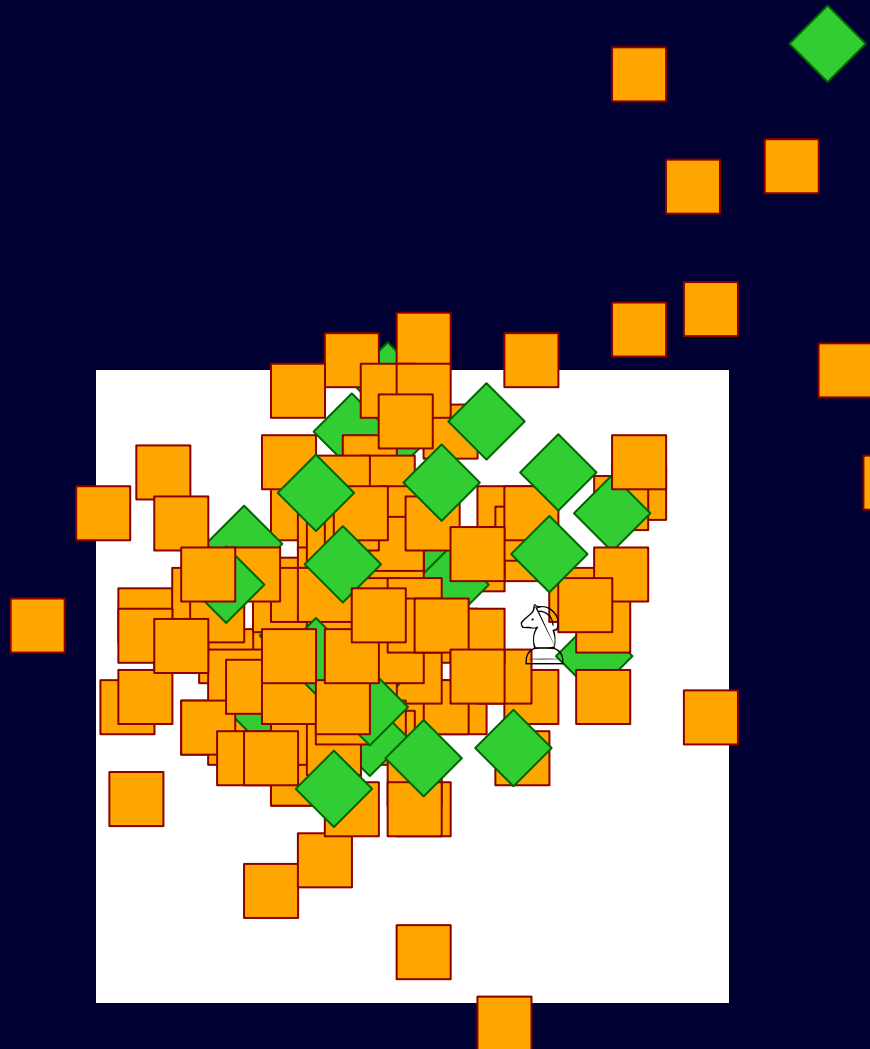


Hans Rudolf Straub

Wie die Künstliche Intelligenz zur Intelligenz kommt

Einblick in eine aktuelle Technologie



Hans Rudolf Straub

Wie die Künstliche Intelligenz zur Intelligenz kommt

Hans Rudolf Straub

Wie die Künstliche Intelligenz zur Intelligenz kommt

Einblick in eine aktuelle Technologie

Z I M – Verlag
St. Gallen



Angaben zum Autor

Hans Rudolf Straub

Email: hrstraub@bluewin.chInternet: <https://hrstraub.ch>

2. Auflage

© 2021 Hans Rudolf Straub

ZIM-Verlag, St. Gallen

<https://fischer-zim.ch/verlag/>

Layout: Wolfram Fischer, Hans Rudolf Straub

Grafik auf Umschlag: Wolfram Fischer

Grafiken: Hans Rudolf Straub

Korrektur: Edith Keim-Vogel

Herstellung: BoD – Books on Demand, Norderstedt

ISBN 978-3-905764-11-6

22224

Vorwort

Künstliche Intelligenz (KI) ist heute ein breit diskutiertes Thema, nicht zuletzt wegen der Erfolge, die sich mit einer aktuellen KI-Form, den neuronalen Netzen, sehr eindrücklich erzielen lassen. Gleichzeitig werden auch die Grenzen und Gefahren klar, z. B. die verheerende Blasenbildung in Internetdiskussionen, die eng mit der Wirkweise der neuronalen Netze zu tun haben und gegen die kein Kraut gewachsen zu sein scheint.

Des Weiteren sind Fragen offen, die an unsere menschliche Existenz rühren: Kann eine Maschine eine eigene Intelligenz haben? Wie verändert KI unsere Überlebensmöglichkeiten auf unserem Planeten? Wie verändert KI die Art und Weise, wie wir denken, z. B. in der Wissenschaft? Und wie ist der Einfluss von KI auf die Weise, wie sich das Leben in unserer Gesellschaft organisiert?

Dabei ist KI nicht neu: Schon ganz am Anfang unseres Umgangs mit Computern haben sich Algorithmus-Pioniere wie Alan Turing mit der Frage der Intelligenz und insbesondere der maschinellen Intelligenz beschäftigt. Und neuronale Netze gibt es seit vielen Jahrzehnten.

KI ist also überhaupt nicht neu. Ich selber baue seit 1989 Programme, die sich mit NLP (Natural Language Processing) beschäftigen, also mit der Interpretation von frei formulierten Texten durch Computer. Sind die so geschaffenen Programme intelligent, wenn sie von Menschen frei erstellte Texte selbstständig «verstehen»?

Zwischen Oktober 2019 und August 2020 habe ich im Internet eine Reihe von Blogbeiträgen zum Thema veröffentlicht, die in der vorliegenden Publikation zusammengefasst sind. Die Texte erscheinen so, wie sie als Blogbeiträge publiziert und weiterhin im Internet lesbar sind.¹

In der Beitragsreihe habe ich die grundsätzlichen Prinzipien der verschiedenen KI-Formen sowie ihre Geschichte, Vorteile und Schwächen dargestellt.

Wie kommen neuronale Netze zu ihrem Wissen? Das entscheidende Grundprinzip, auf dem neuronale Netze beruhen, nämlich der bewertete Lernkorpus, wird erklärt und dabei aufgezeigt, wie und von wem die Bewertung des Korpus durchgeführt wird. Dies führt uns zum Einsatz in Suchmaschinen und zur Blasenbildung im Internet.

Die Bewertung der Daten spielt auch eine wichtige Rolle bei den Programmen des «Deep Learning», die Schach und Go spielen können und ihre Züge dabei selber bewerten. Warum geht das und wo bleibt hier der Unterschied zur menschlichen Form der Intelligenz?

¹ <https://hrstraub.ch/kuenstliche-intelligenz-uebersichtsseite>

Hans Rudolf Straub

Fazit	31
6 Vergleich der Entwicklung der beiden KI-Methoden	33
Zwei KI-Methoden und ihre Herausforderungen	33
Was wurde seit den 90er-Jahren verbessert?	33
Verbreitung der KI-Methoden im Verlauf der Zeit	34
7 Regelbasierte KI: Wo steckt die Intelligenz?	35
Zwei KI-Varianten: regelbasiert und korpusbasiert	35
Aufbau eines regelbasierten Systems	35
Wo sitzt nun die Intelligenz?	36
8 Korpusbasierte KI: Wo steckt die Intelligenz?	37
Vorbemerkung	37
Schritt 1: Erstellung der Datensammlung	37
Schritt 2: Bewertung des Korpus	38
Schritt 3: Training des neuronalen Netzes	40
Schritt 4: Anwendung	41
Fazit	42
9 Was der Korpus weiss – und was nicht	43
Die Erstellung des Korpus	43
Der Zufall regiert im zu kleinen Korpus	44
Raupen- oder Radpanzer?	44
Fazit	45
10 Die Intelligenz in der Suchmaschine	47
Wie kommt die Intelligenz in die Suchmaschine?	47
Trick 1: Lass die Kunden den Korpus trainieren	47
Trick 2: Bewerte die Kunden dabei mit	48
Konsequenzen	48
11 Wie real ist das Wahrscheinliche?	51
Was nicht im Korpus ist, ist für die KI unsichtbar	51
Das neuronale Netz bewertet nach Wahrscheinlichkeit	51
Ausgangslage	51
Beispiel	52
Ist das aber sicher so?	52
Fazit für unsere Suchmaschine	53
Fazit für die korpusbasierte KI generell	53
12 Spiele und Intelligenz (1): Jassen und Schach	55
Schach oder Jassen, was erfordert mehr Intelligenz?	55

Gemeinsamkeiten	55
(a) Klares Spielfeld	55
(b) Klare Spielregeln	55
(c) Klarer Spielverlauf (Zeitverlauf)	56
(d) Klares Spielziel	56
Unterschiede	56
(e) Eindeutige Ausgangssituation?	56
(f) Verdeckte Informationen?	56
(g) Wahrscheinlichkeiten und Emotionen (Psychologie)	57
(h) Kommunikation	57
(i) Der legale Graubereich	58
Fazit	59
13 Spiele und Intelligenz (2): Deep Learning	61
Go und Schach	61
Go und Deep Learning	62
Regel- oder korpusbasiert – oder ein neues, drittes System?	62
Deep Learning ist korpusbasiert	62
Die Bewertung der Korpuseinträge	63
Natürliche Grenzen des Deep Learning	63
1. Ein geschlossenes System	63
2. Ein klar definiertes Ziel	64
Fazit	65
14 Übersicht über die KI-Systeme	67
A: Regelbasierte Systeme	67
A1: Einfacher Automat (Typ Taschenrechner)	67
A2: Wissensbasiertes System	68
B: Korpusbasierte Systeme	70
Drei Untertypen der korpusbasierten KI	72
B1: Typ Mustererkennung	72
B2: Typ Suchmaschine	72
B3: Typ Deep Learning	73
C: Hybride Systeme	73
15 Wo in der Künstlichen Intelligenz steckt nun die Intelligenz?	75
(a) Regelbasierte Systeme	75
(b) Konventionelle korpusbasierte Systeme (Mustererkennung)	75
(c) Suchmaschinen	75
(d) Spielprogramme (Schach, Go, usw.) / Deep Learning	76
Fazit	77

16 Künstliche und natürliche Intelligenz: Der Unterschied	79
Was ist wirkliche Intelligenz?	79
Schach und Go sind geschlossene Systeme	79
Mustererkennung: Offenes oder geschlossenes System?	80
Gibt es Intelligenz ohne Absicht?	81
17 Nachwort	83
KI umfasst mehr als nur neuronale Netze	83
Neuronale Netze sind potent	83
Korpus- oder regelbasiert?	83
Die Möglichkeiten der neuronalen Netze sind beschränkt	84
Intransparenz	84
Kleiner Differenzierungsgrad	84
Was hat die biologische der künstlichen Intelligenz voraus? . .	85
Das Unwahrscheinliche einbeziehen	85
Detaillierter differenzieren	85
Transparenz suchen	85
Kontext bewusst wählen	85
Zielorientiert denken	86
Fazit in einem Satz	86

Abbildungen

1	Aufbau einer korpusbasierten KI	18
2	Suchanfrage mit noch nicht klassifiziertem Panzer an das neuronale Netz	19
3	Beispiel eines multifokalen Begriffsmoleküls	29
4	Schätzung der Verbreitung der KI-Methoden	34
5	Semfinder-System (Überblick)	35
6	Korpusbasiertes System	38
7	Korpus mit Bewertungen (e = eigen, f = fremd)	38
8	Bewertung des Korpus	39
9	Das neuronale Netz lernt das Wissen des Korpus.	40
10	Anwendung eines neuronalen Netzes	41
11	Das neuronale Netz holt das Wissen aus seinem Korpus.	43
12	Erstellung der Zuordnungen im Korpus	43
13	Die Bewertung entscheidet – Korpus mit eigenen und fremden Panzern und entsprechend programmiertem Netz	44
14	Einfacher Automat	67
15	Erstellen einer Wissensbasis	68
16	Anwenden eines wissensbasierten Systems	69
17	Erstellen eines korpusbasierten Systems	70
18	Anwenden eines korpusbasierten Systems	71
19	Die drei Typen der korpusbasierten Systeme	72

1 Zur KI: Schnaps und Panzer

KI im letzten Jahrhundert

KI ist heute ein grosses Schlagwort, war aber bereits in den 80er- und 90er-Jahren des letzten Jahrhunderts ein Thema, das für mich auf meinem Gebiet des Natural Language Processing interessant war. Es gab damals **zwei Methoden**, die gelegentlich als KI bezeichnet wurden und die unterschiedlicher nicht hätten sein können. Das Spannende daran ist, dass diese beiden unterschiedlichen Methoden heute noch existieren und sich weiterhin essenziell voneinander unterscheiden.

KI-1: Schnaps

Die erste, d. h. die Methode, die bereits die allerersten Computerpioniere verwendeten, war eine rein algorithmische, d. h. eine **regelbasierte**. Beispielfür diese Art Regelsysteme sind die *Syllogismen* des Aristoteles:

Prämisse 1: Alle Menschen sind sterblich.

Prämisse 2: Sokrates ist ein Mensch.

Schlussfolgerung: Sokrates ist sterblich.

Der Experte gibt Prämisse 1 und 2 ein, und das System zieht dann selbstständig die Schlussfolgerung. Solche Systeme lassen sich mathematisch untermauern. Mengenlehre und **First-Order-Logic** (Aussagenlogik ersten Grades) gelten oft als sichere mathematische Grundlage. Theoretisch waren diese Systeme somit wasserdicht abgesichert. In der Praxis sah die Geschichte allerdings etwas anders aus.

Probleme ergaben sich durch die Tatsache, dass auch die kleinsten Details in das Regelsystem aufgenommen werden mussten, da sonst das Gesamtsystem «abstürzte», d. h. total abstruse Schlüsse zog. Die Korrektur dieser Details nahm mit der Grösse des abgedeckten Wissens überproportional zu. Die Systeme funktionierten allenfalls für kleine Spezialgebiete, für die klare Regeln gefunden werden konnten, für ausgedehntere Gebiete wurden die Regelbasen aber zu gross und waren nicht mehr wartbar.

Ein weiteres gravierendes Problem war die **Unschärfe**, die vielen Ausdrücken eigen ist, und die mit solchen hart-kodierten Systemen in den Griff zu bekommen ist.

Diese Art KI geriet also zunehmend in die Kritik. Kolportiert wurde z. B. folgender **Übersetzungsversuch**: Ein NLP-Programm übersetzte Sätze vom Englischen ins Russische und wieder zurück, dabei ergab die Eingabe: «*Das Fleisch ist willig, aber der Geist ist schwach*», die Übersetzung: «*Das Steak ist kräftig, aber der Schnaps ist lahm*.»

Die Geschichte hat sich vermutlich nicht genau so zugetragen, aber das Beispiel zeigt die Schwierigkeiten, wenn man versucht, Sprache mit regelbasierten Systemen einzufangen. Die Anfangseuphorie, die seit den 50er-Jahren mit dem «Elektronenhirn» und seiner «maschinellen Intelligenz» verbunden worden war, verblasste, der Ausdruck «Künstliche Intelligenz» wurde obsolet und durch den Ausdruck «Expertensystem» ersetzt, der weniger hochgestochen klang.

Später, d. h. um 2000, gewannen die Anhänger der regelbasierten KI allerdings wieder Auftrieb. Tim Berners-Lee, Pionier des WWW, lancierte zur besseren Benutzbarkeit des Internets die Initiative **Semantic Web**. Die Experten der regelbasierten KI, ausgebildet an den besten technischen Hochschulen der Welt, waren gern bereit, ihm dafür Wissensbasen zu bauen, die sie nun **Ontologien** nannten. Bei allem Respekt vor Berners-Lee und seinem Bestreben, Semantik ins Netz zu bringen, muss festgestellt werden, dass die Initiative Semantic Web nach bald 20 Jahren das Internet nicht wesentlich verändert hat. Meines Erachtens gibt es gute Gründe dafür: Die Methoden der klassischen mathematischen Logik sind zu rigid, die komplexen Vorgänge des Denkens nachzuvollziehen – mehr dazu in meinen anderen Beiträgen, insbesondere zur statischen und dynamischen Logik.¹ Jedenfalls haben weder die klassischen regelbasierten Expertensysteme des 20. Jahrhunderts noch die Initiative «Semantic Web» die hoch gesteckten Erwartungen erfüllt.

↑ S. 30: Punkt 3

KI-2: Panzer

In den 90er-Jahren gab es aber durchaus auch schon Alternativen, die versuchten, die Schwächen der rigiden Aussagenlogik zu korrigieren. Dazu wurde das mathematische Instrumentarium erweitert.

Ein solcher Versuch war die Fuzzy Logic. Eine Aussage oder eine Schlussfolgerung war nun nicht mehr eindeutig wahr oder falsch, sondern der Wahrheitsgehalt konnte gewichtet werden. Neben Mengenlehre und Prädikatenlogik hielt nun auch die Wahrscheinlichkeitstheorie Einzug ins mathematische Instrumentarium der Expertensysteme. Doch einige Probleme blieben: Wieder musste genau und aufwendig beschrieben werden,

¹ <https://hrstraub.ch/logodynamik>

welche Regeln gelten. Die Fuzzy Logic gehört also ebenfalls zur regelbasierten KI, wenn auch mit Wahrscheinlichkeiten versehen. Heute funktionieren solche Programme in kleinen, wohlabgegrenzten technischen Nischen perfekt, haben aber darüberhinaus keine Bedeutung.

Eine andere Alternative waren damals die **Neuronalen Netze**. Sie galten als interessant, allerdings wurden ihre praktischen Anwendungen eher etwas belächelt. Folgende Geschichte wurde dazu herumgereicht:

Die amerikanische Armee – seit jeher ein wesentlicher Treiber der Computertechnologie – soll ein neuronales Netz zur Erkennung von eigenen und fremden Panzern gebaut haben. Ein neuronales Netz funktioniert so, dass die Schlussfolgerungen über mehrere Schichten von Folgerungen vom System selber gefunden werden. Der Mensch muss also keine Regeln mehr eingeben, diese werden vom System selber erstellt.

Wie kann das System das? Es braucht dazu einen **Lernkorpus**. Bei der Panzererkennung war das eine Serie von Fotos von amerikanischen und russischen Panzern. Für jedes Foto war also bekannt, ob amerikanisch oder russisch, und das System wurde nun so lange trainiert, bis es die geforderten Zuordnungen selbstständig erstellen konnte. Die Experten nahmen auf das Programm nur indirekt Einfluss, indem sie den Lernkorpus aufbauten; das Programm stellte die Folgerungen im neuronalen Netz selbstständig zusammen – ohne dass die Experten genau wussten, aus welchen Details das System mit welchen Regeln welche Schlüsse zog. Nur das Resultat musste natürlich stimmen. Wenn das System nun den Lernkorpus vollkommen integriert hatte, konnte man es testen, indem man ihm einen neuen Input zeigte, z. B. ein neues Panzerfoto, und es wurde erwartet, dass es mit den aus dem Lernkorpus gefundenen Regeln das neue Bild korrekt zuordnete. Die Zuordnung geschah, wie gesagt, selbstständig durch das System, ohne dass der Experte weiteren Einfluss nahm und ohne dass er genau wusste, wie im konkreten Fall die Schlüsse gezogen wurden.

Das funktionierte, so wurde erzählt, bei dem Panzererkennungsprogramm perfekt. So viele Fotos dem Programm auch gezeigt wurden, stets erfolgte die korrekte Zuordnung. Die Experten konnten selber kaum glauben, dass sie wirklich ein Programm mit einer hundertprozentigen Erkennungsrate erstellt hatten. Wie konnte so etwas sein? Schliesslich fanden sie den Grund: Die Fotos der amerikanischen Panzer waren in Farbe, diejenigen der russischen schwarzweiss. Das Programm musste also nur die Farbe erkennen, die Silhouetten der Panzer waren irrelevant.

Regelbasiert versus korpusbasiert

Die beiden Anekdoten zeigen, welche Probleme damals auf die regelbasierte und die korpusbasierte KI warteten.

Bei der **regelbasierten KI** waren es:

- die *Rigidität* der mathematischen Logik,
- die *Unschärfe* unserer Wörter,
- die Notwendigkeit, sehr grosse *Wissensbasen* aufzubauen,
- die Notwendigkeit, *Fachexperten* für die Wissensbasen einzusetzen.

Bei der **korpusbasierten KI** waren es:

- die *Intransparenz* der Schlussfolgerungs-Wege,
- die Notwendigkeit, einen sehr grossen und relevanten *Lernkorpus* aufzubauen.

Ich hoffe, dass ich mit den beiden oben beschriebenen, zugegebenermassen etwas unfairen Beispielen den Charakter und die Wirkungsweise der beiden KI-Typen habe darstellen können, mitsamt den Schwächen, die die beiden Typen jeweils kennzeichnen.

Die Herausforderungen bestehen selbstverständlich weiterhin. Im Folgenden werde ich darstellen, wie die beiden KI-Typen darauf reagiert haben und wo bei den beiden Systemen nun wirklich die Intelligenz sitzt. Als Erstes schauen wir die korpusbasierte KI an.

2 Die korpusbasierte KI überwindet ihre Schwächen

Im vorangegangenen Kapitel erwähnte ich die beiden prinzipiellen Herangehensweisen, mit denen versucht wird, dem Computer Intelligenz beizubringen, nämlich die *regelbasierte* und die *korpusbasierte*. Bei der regelbasierten steckt die Intelligenz in einem Regelpool, der von Menschen bewusst konstruiert wird. Bei der korpusbasierten Methode steckt das Wissen im Korpus, d. h. in einer Datensammlung, welche von einem raffinierten Programm analysiert wird.

Beide Methoden haben ihre Leistungen seit den 90er-Jahren gewaltig steigern können. Am eindrücklichsten ist dies bei der korpusbasierten Methode geschehen, die heute als eigentliche Künstliche Intelligenz gilt und in der breiten Öffentlichkeit für Schlagzeilen sorgt. Worauf beruhen die entscheidenden Verbesserungen der beiden Methoden? – Ich werde gleich auf beide Methoden und ihre Verbesserungen eingehen. Als Erstes sehen wir uns an, wie die korpusbasierte KI funktioniert.

Wie funktioniert die korpusbasierte KI?

Eine korpusbasierte KI besteht aus zwei Teilen:

1. Korpus
2. Algorithmen (neuronales Netz)

Der Korpus, auch Lernkorpus genannt, ist eine Sammlung von Daten. Dies können z. B. Fotografien von Panzern oder Gesichtern sein, aber auch Sammlungen von Suchanfragen, z. B. von Google. Wichtig ist, dass der Korpus die Daten bereits bewertet enthält. Im Panzerbeispiel ist im Korpus vermerkt, ob es sich um eigene oder feindliche Panzer handelt. In der Gesichtersammlung ist vermerkt, um wessen Gesicht es sich jeweils handelt; bei den Suchanfragen speichert Google, welcher Link der Suchende anklickt, d. h. welcher Vorschlag von Google erfolgreich ist. Im Lernkorpus steckt also das Wissen, das die korpusbasierte KI verwenden wird.

Nun muss die KI lernen. Das Ziel ist, dass die KI ein neues Panzerbild, ein neues Gesicht oder eine neue Suchanfrage korrekt zuordnen kann. Dazu verwendet die KI das Wissen im Korpus, also z. B. die Bilder der Panzersammlung, wobei bei jedem Bild vermerkt ist, ob es sich um eigene oder fremde Panzer handelt – in **Abb. 1** dargestellt durch die kleinen grauen und grünen Etiketten links von jedem Bild. Diese Bewertungen sind ein notwendiger Teil des Korpus.

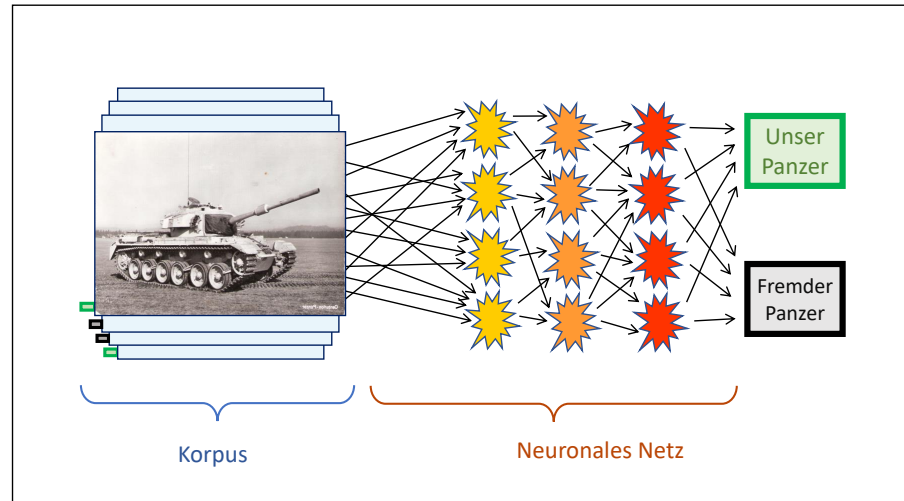


Abb. 1: Aufbau einer korpusbasierten KI

Jetzt kommt der zweite Bestandteil der korpusbasierten KI ins Spiel, der Algorithmus. Im Wesentlichen handelt es sich um ein neuronales Netz. Es besteht aus mehreren Schichten von «Neuronen», die Inputsignale aufnehmen, gegeneinander verrechnen und dann ihre eigenen Signale an die nächsthöhere Schicht ausgeben. In **Abb. 1** ist dargestellt, wie die erste (gelbe) Neuronenschicht die Signale (Pixel) aus dem Bild aufnimmt und nach einer Verrechnung dieser Signale eigene Signale an die nächste (orange) Schicht weitergibt, bis am Schluss das Netz zum Resultat «eigener» oder «fremder» Panzer gelangt. Die Verrechnungen (Algorithmen) der Neuronen werden beim Training so lange verändert und angepasst, bis das Gesamtnetz bei jedem Bild das korrekte Resultat liefert.

Wenn jetzt ein neues, noch unbewertetes Bild dem neuronalen Netz vorgelegt wird, verhält sich dieses genau gleich wie bei den anderen Bildern. Wenn das Netz gut trainiert worden ist, sollte der Panzer vom Programm selbstständig zugeordnet werden können, d. h. das neuronale Netz erkennt, ob das Bild einen eigenen oder fremden Panzer darstellt (**Abb. 2**).

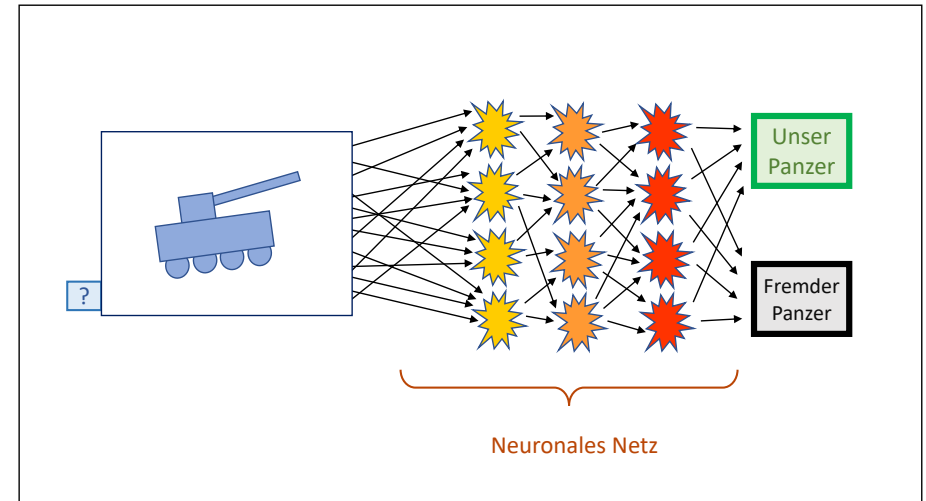


Abb. 2: Suchanfrage mit noch nicht klassifiziertem Panzer an das neuronale Netz

Die Bedeutung des Datenkorpus für die korpusbasierte KI

Die korpusbasierte KI findet ihr Detailwissen im eigens für sie bereitgestellten Korpus vor und wertet die Verbindungen aus, die sie dort antrifft. Der Korpus enthält somit das Wissen, welches die korpusbasierte KI auswertet. Das Wissen besteht in unserem Beispiel in der *Verbindung* der Fotografie – also einer Menge von wild angeordneten Pixeln – mit einer einfachen binären Information (unser Panzer/fremder Panzer). Dieses Wissen findet sich im Korpus bereits *bevor* eine Auswertung durch die Algorithmen stattfindet. Die Algorithmen der korpusbasierte KI finden also nichts heraus, was nicht im Korpus steckt. Allerdings: Das im Korpus gefundene Wissen kann die korpusbasierte KI nun auch auf neue und noch nicht bewertete Fälle anwenden.

2 Die korpusbasierte KI überwindet ihre Schwächen

Die Herausforderungen an die korpusbasierte KI

Die Herausforderungen an die korpusbasierte KI sind eindeutig:

1. **Grösse des Korpus:** Je mehr Bilder sich im Korpus befinden, umso sicherer kann die Zuordnung erfolgen. Ein zu kleiner Korpus bringt Fehlresultate. Die Grösse des Korpus ist für die **Präzision** und Zuverlässigkeit der Resultate entscheidend.
2. **Hardware:** Die Rechenleistung, welche die korpusbasierte KI benötigt, ist sehr gross; und sie wird umso grösser, je präziser die Methode sein soll. Die Performance der Hardware entscheidet über die praktische Anwendbarkeit der Methode.

Dadurch wird schnell klar, wie die korpusbasierte KI ihre Leistung in den letzten zwei Jahrzehnten so eindrücklich verbessern konnte:

1. Die Datenmengen, welche Google und andere Organisationen im Internet sammeln können, sind drastisch angestiegen. Google profitiert dabei von einem nicht unbedeutenden Verstärkungseffekt: Je mehr Anfragen Google bekommt, umso besser wird der Korpus und damit seine Trefferquote. Je besser die Trefferquote, umso mehr Anfragen bekommt Google.
2. Die Hardware, welche zur Auswertung der Daten benötigt wird, wird immer günstiger und performanter. Internetfirmen und andere Organisationen verfügen heute über riesige Serverfarmen, welche die rechenintensiven Auswertungen der korpusbasierten KI erst möglich machen.

Neben dem Korpus und der Hardware spielt natürlich auch die Raffinesse der Algorithmen eine Rolle. Die Algorithmen waren aber auch schon vor Jahrzehnten nicht schlecht. Im Vergleich zu den beiden anderen Faktoren – Hardware und Korpus – spielt der Fortschritt bei den Algorithmen für den beeindruckenden Erfolg der korpusbasierten KI nur eine bescheidene Rolle.

Der Erfolg der korpusbasierten KI

Die Herausforderungen an die korpusbasierte KI wurden von den grossen Firmen und Organisationen äusserst erfolgreich angegangen.

Auf Basis der oben erfolgten Beschreibung der Funktionsweise sollten aber auch die systemimmanenten und in den Medien etwas weniger prominent platzierten Schwächen der korpusbasierten KI erkennbar werden. In Kapitel 9 werde ich genauer darauf eingehen.

↑ Kapitel 9: S. 43

3 Die Herausforderungen an die regelbasierte KI

Regelbasiert im Vergleich zu korpusbasiert

Die korpusbasierte KI (Typus «Panzer») konnte ihre Schwächen erfolgreich überwinden (siehe vorangegangenes Kapitel). Dafür reichte eine Kombination von «Brute Force» (verbesserte Hardware) und einem idealen Opportunitätsfenster: Als nämlich während der superheissen Expansionsphase des Internets Firmen wie Google, Amazon, Facebook und viele andere grosse Datenmengen sammeln und damit ihre Datenkorpora füttern konnten. Und mit einem ausreichend grossen Datenkorpus steht und fällt die korpusbasierte KI.

Für die regelbasierte KI aber reichte «Brute Force» nicht aus. Es nützte auch nichts, viele Daten zu sammeln, da für den Regelbau die Daten auch organisiert werden müssen – und zwar grossenteils von Hand, also durch menschliche Fachexperten.

Herausforderung 1: Unterschiedliche Mentalitäten

Nicht alle Menschen sind gleichermassen davon fasziniert, Algorithmen zu bauen. Es braucht dazu eine besondere Art Abstraktionsfähigkeit, gepaart mit einer sehr gewissenhaften Ader – jedenfalls was die Abstraktionen betrifft. Jeder noch so kleine Fehler im Regelbau wird sich unweigerlich auswirken. Mathematiker verfügen sehr ausgeprägt über diese hier gefragte konsequent-gewissenhafte Mentalität, aber auch Naturwissenschaftler und Ingenieure zeichnen sich vorteilhaft dadurch aus. Natürlich müssen auch Buchhalter gewissenhaft sein, für den Regelbau der KI ist aber zusätzlich noch Kreativität gefragt.

Verkäufer, Künstler und Ärzte hingegen arbeiten in anderen Bereichen. Oft ist Abstraktion eher nebensächlich, und das Konkrete ist wichtig. Auch das Einfühlungsvermögen in andere Menschen kann sehr wichtig sein. Oder man muss schnell und präzise handeln können, z. B. als Chirurg. Diese Eigenschaften sind alle sehr wertvoll, für den Algorithmenbau aber weniger wichtig.

Das ist für die regelbasierte KI ein Problem. Denn für den Regelbau braucht es sowohl die Fähigkeiten des einen als auch das Wissen des anderen Lagers: Es braucht die Mentalität, die einen guten Algorithmenmischer ausmacht, gepaart mit der Denkweise und dem Wissen des Fachgebiets,

↑ S. 14

auf das sich die Regeln beziehen. Eine solche Kombination des Fachgebietswissens mit dem Talent zur Abstraktion ist selten zu finden. In den Krankenhäusern, in denen ich gearbeitet habe, waren die beiden Kulturen in ihrer Getrenntheit ganz klar ersichtlich. Hier die Ärzte, die Computer höchstens für die Rechnungsstellung oder für gewisse teure technische Apparate akzeptierten, die Informatik allgemein aber geringschätzten, und dort die Informatiker, die keine Ahnung davon hatten, was die Ärzte taten und wovon sie überhaupt sprachen. Die beiden Lager gingen sich meist einfach aus dem Weg. Selbstverständlich war es da nicht verwunderlich, dass die für die Medizin gebauten Expertensysteme meist nur für ganz kleine Teilgebiete funktionierten, wenn sie nicht im blossen Experimentierstadium verharrten.

Herausforderung 2: Wo finde ich die Experten?

Experten, die kreativ und in den beiden Mentalitätslagern gleichermaßen zuhause sind, sind selbstverständlich schwer zu finden. Erschwerend kommt hinzu: Es gibt kaum Ausbildungsstätten für diese Art Experten. Realistisch sind auch folgende Fragen: Wo sind die Auszubildenden, die sich mit den aktuellen Herausforderungen auskennen? Welche Diplome gelten wofür? Und wie evaluiert ein Geldgeber auf diesem neuen Gebiet, ob die eingesetzten Experten taugen und die Projektrichtung stimmt?

Herausforderung 3: Schiere Menge an nötigen Detailregeln

Dass eine grosse Menge an Detailwissen nötig ist, um in einer Realsituation sinnvolle Schlüsse zu ziehen, war schon für die korpusbasierte KI eine Herausforderung. Denn erst mit wirklich grossen Korpora, d. h. dank des Internets und gesteigerter Computerleistung gelang es ihr, die riesige Menge an Detailwissen zu erfassen, das für jedes realistische Expertensystem eine der Basisvoraussetzungen ist.

Für die regelbasierte KI ist es aber besonders schwierig, die grosse Wissensmenge bereitzustellen, denn sie braucht für die Wissenserstellung Menschen, welche die grosse Wissensmenge von Hand in computergängige Regeln fassen. Das ist eine sehr zeitraubende Arbeit, die zudem die schwierig zu findenden menschlichen Fachexperten erfordert, die den oben genannten Herausforderung 1 und 2 genügen.

In dieser Situation stellt sich die Frage, wie grössere und funktionierende Regelsysteme überhaupt gebaut werden können? Gibt es eventuell Möglichkeiten, den Bau der Regelsysteme zu vereinfachen?

Herausforderung 4: Komplexität

Wer je versucht hat, ein Fachgebiet wirklich mit Regeln zu unterfüttern, merkt, dass er schnell an komplexe Fragen stösst, für die er in der Literatur keine Lösungen findet. In meinem Gebiet des Natural Language Processing (NLP) ist das offensichtlich. Die Komplexität ist hier nicht zu übersehen. Deshalb muss unbedingt auf sie eingegangen werden. Mit anderen Worten: Das Prinzip Hoffnung reicht nicht, sondern die Komplexität muss thematisiert und intensiv studiert werden.

Was Komplexität bedeutet, und wie man ihr begegnen kann, darauf möchte ich in einem weiteren Beitrag eingehen. Selbstverständlich darf dabei die Komplexität nicht zu einer übermässigen Regelvermehrung führen (siehe Herausforderung 3). Die Frage, die sich für die regelbasierte KI stellt, ist deshalb: Wie kann ein Regelsystem gebaut werden, das Detailhaltigkeit und Komplexität berücksichtigt, dabei aber einfach und übersichtlich bleibt?

Die gute Botschaft ist: Auf diese Frage gibt es durchaus Antworten.

16 Künstliche und natürliche Intelligenz: Der Unterschied

Wie beim Schachprogramm ist die künstliche Intelligenz nicht imstande, das Ziel selbständig herauszufinden, beim Schachprogramm versteht sich das Ziel (Schachmatt) von selber, bei der Mustererkennung müssen sich die beteiligten Bewerter über das Ziel (fremde/eigene, Rad-/Kettenpanzer) vorgängig einig sein. In beiden Fällen kommen Ziel und Absicht von aus-sen.

Natürliche Intelligenz hingegen kann sich selber darüber klar werden, was wichtig und was unwichtig ist und welche Ziele sie verfolgt. Die aktive Absicht ist m. E. eine unverzichtbare Eigenschaft der natürlichen Intelligenz und kann nicht künstlich konstruiert werden.

Fazit 3:

*Im Gegensatz zur künstlichen zeichnet sich die natürliche Intelligenz dadurch aus, dass sie die **eigenen Absichten beurteilen** und **bewusst ausrichten** kann.*

↑ S. 17

17 Nachwort

KI umfasst mehr als nur neuronale Netze

Neuronale Netze sind potent

Was unter Künstlicher Intelligenz (KI) allgemein verstanden wird, sind sogenannte Neuronale Netze.

Neuronale Netze sind potent und für ihre Anwendungsgebiete unschlagbar. Sie erweitern die technischen Möglichkeiten unserer Zivilisation massgeblich auf vielen Gebieten. Trotzdem sind neuronale Netze nur eine Möglichkeit, «intelligente» Computerprogramme zu organisieren.

Korpus- oder regelbasiert?

Neuronale Netze sind **korpusbasiert**, d. h. ihre Technik basiert auf einer Datensammlung, dem Korpus, der von aussen in einer Lernphase Datum für Datum bewertet wird. Das Programm erkennt anschliessend in der Bewertung der Daten selbständig gewisse Muster, die auch für bisher unbekannte Fälle gelten. Der Prozess ist automatisch, aber auch intransparent. In einem realen Einzelfall ist nicht klar, welche Gründe für die Schlussfolgerungen herangezogen worden sind. Wenn der Korpus aber genügend gross und korrekt bewertet ist, ist die Präzision der Schlüsse ausserordentlich hoch.

Grundsätzlich anders funktionieren **regelbasierte** Systeme. Sie brauchen keine Datensammlung, sondern eine Regelsammlung. Die Regeln werden von Menschen erstellt und sind transparent, d. h. leicht les- und veränderbar. Regelbasierte Systeme funktionieren allerdings nur mit einer adäquaten Logik (dynamische, nicht statische Logik) und einer für komplexe Semantiken geeigneten, multifokalen Begriffsarchitektur; beides wird in den entsprechenden Hochschulinstituten bisher kaum gelehrt.

Aus diesem Grund stehen regelbasierte Systeme heute eher im Hintergrund, und was allgemein unter KI verstanden wird, sind neuronale Netze, also korpusbasierte Systeme.

Die Möglichkeiten der neuronalen Netze sind beschränkt

Der unbezweifelbare Erfolg der neuronalen Netze lässt ihre Schwächen in den Hintergrund treten. Als korpusbasierte Systeme sind neuronale Netze völlig von der vorgängig erfolgten Datensammlung, dem Korpus abhängig. Prinzipiell kann nur die Information, die auch im Korpus steckt, vom neuronalen Netz überhaupt gesehen werden. Zudem muss der Korpus bewertet werden, was durch menschliche Experten erfolgt. Was nicht im Korpus steckt, befindet sich ausserhalb des Horizonts des neuronalen Netzes. Fehler im Korpus oder in seiner Bewertung führen zu Fehlern im neuronalen Netz:

↑ S. 44 und S. 51

↑ S. 45

Intransparenz

Welche Datensätze im Korpus zu welchen Schlüssen im neuronalen Netz führen, lässt sich im Nachhinein nicht rekonstruieren. Somit können Fehler im neuronalen Netz nur mit beträchtlichem Aufwand korrigiert werden.

Andererseits ist es auch nicht nötig, die Schlussfolgerungen wirklich zu verstehen. So kann ein neuronales Netz ein Lächeln auf einem fotografierten Gesicht sehr gut erkennen, obwohl wir nicht bewusst angeben können, welche Pixelkombinationen nun genau für das Lächeln verantwortlich sind.

Kleiner Differenzierungsgrad

Wie viele Ergebnismöglichkeiten (Outcomes) kann ein neuronales Netz unterscheiden? Jedes mögliche Ergebnis muss in der Lernphase einzeln geschult und abgegrenzt werden. Dafür müssen genügend Fälle im Korpus vorhanden sein. Der Aufwand bezüglich Korpusgrösse steigt dabei nicht proportional, sondern exponentiell. Dies führt dazu, dass Fragen mit wenig Ergebnismöglichkeiten von neuronalen Netzen sehr gut, solche mit vielen unterschiedlichen Antworten nur mit überproportionalem Aufwand gelöst werden können. Am besten eignen sich deshalb binäre Antworten, z. B.: Ist der Twittertext von einer Frau oder einem Mann geschrieben? Zuweisungen mit vielen Outcome-Möglichkeiten hingegen eignen sich schlecht.

↑ S. 44

Was hat die biologische der künstlichen Intelligenz voraus?

Das Unwahrscheinliche einbeziehen

Neuronale Netze bewerten die Wahrscheinlichkeit eines Ergebnisses. Dies entspricht einem sehr flachen Denkvorgang, denn nicht nur das Wahrscheinliche ist möglich. Gerade eine unwahrscheinliche Wendung kann ganz neue Perspektiven öffnen, im Leben wie im Denken. Das automatische Vorgehen der neuronalen Netze aber ist ein Denkkorsett, das stets das Wahrscheinliche erzwingt.

↑ S. 51

Detaillierter differenzieren

Neuronale Netze werden umso unpräziser, je mehr Details sie unterscheiden sollen. Schon mit wenigen Resultatmöglichkeiten (Outcomes) sind sie überfordert. Biologische Intelligenz hingegen kann sich je nach Fragestellung in sehr differenzierte Ergebniswelten eindenken, mit einer Vielfalt von Ergebnismöglichkeiten.

↑ S. 44

Transparenz suchen

Weshalb komme ich im Denken zu einem bestimmten Ergebnis? Wie ist der Denkverlauf? Nur wenn ich mein Denken hinterfragen kann, kann ich es verbessern. Neuronale Netze hingegen können ihre Schlüsse nicht hinterfragen. Sie folgern einfach das, was der Korpus und dessen von aussen erfolgte Bewertung ihnen vorgeben.

Kontext bewusst wählen

Je nach Fragestellung wählt das menschliche Denken einen Kontext mit entsprechenden Musterbeispielen und bereits erkannten Regeln. Diese Auswahl ist **aktiv** und stellt das bewusste Moment im Denken dar: Worüber denke ich nach?

↑ S. 40

Die Auswahl des Kontexts entscheidet natürlich auch über die möglichen Outcomes und gültigen Regelmuster. Eine aktive Bewertung und Filterung des Kontexts erbringt die biologische Intelligenz automatisch, sie liegt jedoch ausserhalb der Möglichkeiten eines neuronalen Netzes.

Zielorientiert denken

Was habe ich für Ziele? Was ist für mich wichtig? Was hat für mich Bedeutung? – Solche Fragen richten mein Denken aus; Aristoteles spricht von der «Causa finalis», dem **Wozu**, d. h. dem Ziel als Beweggrund. Neuronale Netze kennen kein eigenes Ziel, d. h. dieses steckt stets in der vorgängigen Bewertung des Korpus, die von aussen erfolgt. Das Ziel ist nichts, was das neuronale Netz selbstständig oder im Nachhinein verändern kann. Ein neuronales Netz ist stets und vollständig fremdbestimmt.

Eine biologische Intelligenz hingegen kann sich über ihre Ziele Gedanken machen und das Denken entsprechend verändern und ausrichten. Diese Autonomie über die eigenen Ziele zeichnet die biologische Intelligenz aus und ist ein wesentliches Element des Ideals des freien Menschen.

Fazit in einem Satz

*Die künstliche Intelligenz der **neuronalen Netze** ist hochpotent, aber nur für umschriebene, eng begrenzte Fragestellungen einsetzbar und hat mit wirklicher, d. h. aktiver Intelligenz nichts zu tun.*

↑ S. 64 und S. 81

Publikationen von Hans Rudolf Straub im ZIM-Verlag

Hans Rudolf Straub

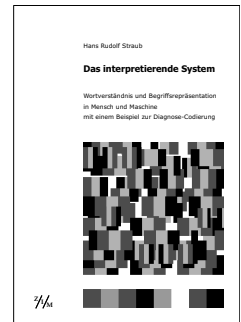
Das interpretierende System

Wortverständnis und Begriffsrepräsentation in Mensch und Maschine mit einem Beispiel zur Diagnose-Codierung

4. Auflage. 2015, 176 S., 13 Tab. und 62 Abb., kt.

€ 22.00 / SFr. 22.00

ISBN 978-3-905764-09-3 (ZIM)



Hans Rudolf Straub

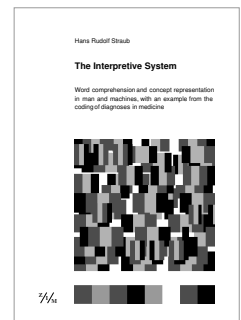
The Interpretive System

Word comprehension and concept representation in man and machines, with an example from the coding of diagnoses in medicine

English Edition. 2020, 157 S., 13 Tab. und 62 Abb., kt.

€ 29.00 / SFr. 29.00

ISBN 978-3-905764-10-9 (ZIM)



<https://fischer-zim.ch/verlag>

Haben Sie sich schon gefragt, woher die Intelligenz in der künstlichen Intelligenz stammt? Die Künstliche Intelligenz, etwa in Suchmaschinen oder Programmen zur Gesichtserkennung ist ja nichts anderes als ein Computerprogramm, das vorgängig mit Daten gefüttert worden ist. Ein raffiniertes und komplexes Programm zwar, aber letztlich nichts als ein Automat, eine Maschine. Und diese Maschine verhält sich intelligent. Wie kann das sein?

Nun, darauf gibt es durchaus eine Antwort. Basierend auf meiner Berufspraxis im Bereich des NLP (Natural Language Processing) und der automatisierten Zuweisung von medizinischen Codes (ICD, CHOP, OPS, SNOMED) zu Freitexten, habe ich eine Serie von Blogbeiträgen im Web publiziert, die nun in Buchform erscheinen.

Die Beiträge stellen verschiedene Formen von KI vor und erklären kurz ihre prinzipiellen Wirkungsweisen. Für jede vorgestellte KI-Form wird gezeigt, an welcher Stelle und wie die benötigte Intelligenz zum Programm hinzu kommt. Selbstverständlich unterscheiden sich Taschenrechner, Suchmaschinen, regelbasierte NLP-Systeme und Deep Learning Systeme (Schach, Go) in dieser Beziehung deutlich.

Was aber sind die Konsequenzen dieser Programme und ihres Einsatzes? Was können sie? Und was bewirken sie? Und was ist der Unterschied zur menschlichen, also zur biologischen Intelligenz?

Z I M – Verlag
Zentrum für Informatik



ISBN 978-3-905764-11-6

