

Wissensbasierte Systeme

Master of Science

Prof. Dr. Rethmann

Fachbereich Elektrotechnik und Informatik
Hochschule Niederrhein

WiSe 2019/20

Künstliche Intelligenz:

- Geschichte und Anwendungen
- Wissensrepräsentation
- Zustandsraumsuche
- Aussagen- und Prädikatenlogik
- Einführung in PROLOG



- S. Russell, P. Norvig:
Artificial Intelligence – A Modern Approach.
Prentice Hall, 2002.
- George F. Luger:
Künstliche Intelligenz – Strategien zur Lösung komplexer Probleme.
Addison-Wesley, 2001.
- C. Beierle, G. Kern-Isberner:
Methoden wissensbasierter Systeme.
Vieweg Verlag, 2003.
- G. Görz, C.-R. Rollinger, J. Schneeberger (Hrsg.):
Handbuch der künstlichen Intelligenz.
4. Auflage, Oldenbourg Verlag, 2003.

- U. Lämmel, J. Cleve:
Lehr- und Übungsbuch Künstliche Intelligenz.
Fachbuchverlag Leipzig, 2001.
- J. Heinsohn, R. Socher-Ambrosius:
Wissensverarbeitung – Eine Einführung.
Spektrum Akademischer Verlag, 1999.
- Uwe Schöning:
Logik für Informatiker.
Spektrum Akademischer Verlag, 2000.
- Ivan Bratko:
PROLOG – Programming for Artificial Intelligence.
3rd edition, Addison-Wesley, 2000.

- I. Gerdes, F. Klawonn, R. Kruse:
Evolutionäre Algorithmen.
Vieweg Verlag, 2004.
- Hans-Jürgen Zimmermann:
Operations Research.
Vieweg Verlag, 2005.
- Andreas Zell:
Simulation Neuronaler Netze.
Addison-Wesley, 1994.

- *Geschichte und Anwendungen*
- Wissensrepräsentation
- Zustandsraumsuche
- Aussagen- und Prädikatenlogik
- Einführung in PROLOG

- *Motivation*
- Turing-Test
- neuronale Netze
- Agenten
- Entscheidungsbäume
- fallbasiertes Schließen

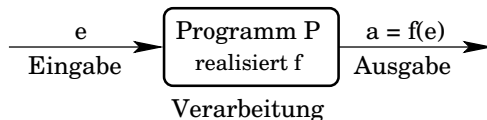
Ein Programm P realisiert eine Funktion f , so dass zu einer Eingabe e eine Ausgabe $a = f(e)$ berechnet wird.

EVA: Eingabe, Verarbeitung, Ausgabe.

Das Programm, und damit die Funktion, wird explizit in einer Programmiersprache wie C# oder Java implementiert.

Jedes Programm berechnet also eine Funktion $f : \mathbb{D} \rightarrow \mathbb{W}$, wobei $\mathbb{D}$ der Definitionsbereich und $\mathbb{W}$ der Wertebereich ist.

Die Funktion bildet ein Argument $z \in \mathbb{D}$ auf einen Wert $f(z) \in \mathbb{W}$ ab. Notation: $z \mapsto f(z)$



Beispiel: Sortiere die gegebenen Zahlen $x_1, \dots, x_n$ aufsteigend.

Dann bildet f eine Eingabe $e = (x_1, \dots, x_n)$ ab auf eine Ausgabe $a = f(e) = (x_{\pi(1)}, \dots, x_{\pi(n)})$, wobei π eine Permutation ist und $x_{\pi(i)} \leq x_{\pi(j)}$ für $1 \leq i < j \leq n$ gilt.

Algorithmen wie Quicksort, Mergesort oder Heapsort realisieren die Funktion f , indem sie beschreiben, wie eine Zahlenfolge sortiert wird.

Was tun wir aber, wenn wir die zu realisierende Funktion f nicht kennen? Wie soll f dann als Programm realisiert werden?

Wie sieht eine Funktion f aus, die bei Eingabe eines Bildes, bestehend aus 32×32 Pixeln mit bis zu 256 Grauwerten pro Pixel, die dort enthaltenen Ziffern erkennt? Es soll also eine Funktion $f : \{0, \dots, 255\}^{32 \times 32} \rightarrow \{0, 1, \dots, 9\}$ berechnet werden.

Was können Menschen besser als Maschinen?

Bilder erkennen, Texte übersetzen, Sprache verarbeiten, Spiele spielen, Planen, Beweisen, Schlussfolgern, Lernen, neue Situationen erfassen und beurteilen, usw.

Menschen lernen durch Beispiele bzw. Analogie und nutzen Ähnlichkeiten: Vorhandenes Wissen wird an neue Situationen angepasst. Ist das Intelligenz?

Voraussetzung ist, dass für ähnliche Situationen bereits Problemlösungen bekannt sind.

Wenn eine neue Situation N ähnlich zu einer anderen, bereits bekannten Situation A ist, und diese Situation A früher mit einem Verfahren V gelöst wurde, unterscheide 2 Fälle:

- War V für A erfolgreich, wende V auf N an,
- sonst wähle ein anderes Verfahren oder passe V an, so dass es auf N angewendet werden kann.

Was können Menschen besser als Maschinen?

Beim fallbasierten Schließen besteht die Wissensbasis nicht aus generischen Regeln, sondern aus einer Sammlung von Fällen, in denen spezifische frühere Erfahrungen gespeichert sind.

Neue Probleme, mit denen das System konfrontiert wird, werden dadurch gelöst, dass der relevanteste Fall aus der Falldatensammlung heraus gesucht wird und dessen Lösung in geeigneter Form auf das neue Problem übertragen wird.

Schließen kann als erinnerungsbasierter Prozess verstanden werden.

Dazu müssen Ähnlichkeiten bestimmt werden, wozu wir wieder Funktionen angeben müssen! Anders bei neuronalen Netzen ...

Was können Menschen besser als Maschinen?

Mittels neuronaler Netze lässt sich das Problem der Ziffernerkennung heute sehr gut lösen, denn wir besitzen heute eine sehr große Anzahl von handgeschriebenen Ziffern.

Im wesentlichen basiert die gute Erkennung darauf, dass sich die Netze „erinnern“, eine ähnliche Ziffer schon einmal während des Trainings gesehen zu haben und daher der Eingabe eine Ziffer zuordnen können.

Je mehr unterschiedliche Schreibweisen während des Trainings verarbeitet wurden, umso höher die Genauigkeit.

Was ist Wissen?

- **Philosophie:** Wissen ist wahrer, gerechtfertigter Glaube.
- **Wissensrepräsentation:** Wissen ist Information über die Welt, ohne Anspruch auf deren Wahrheitsgehalt.

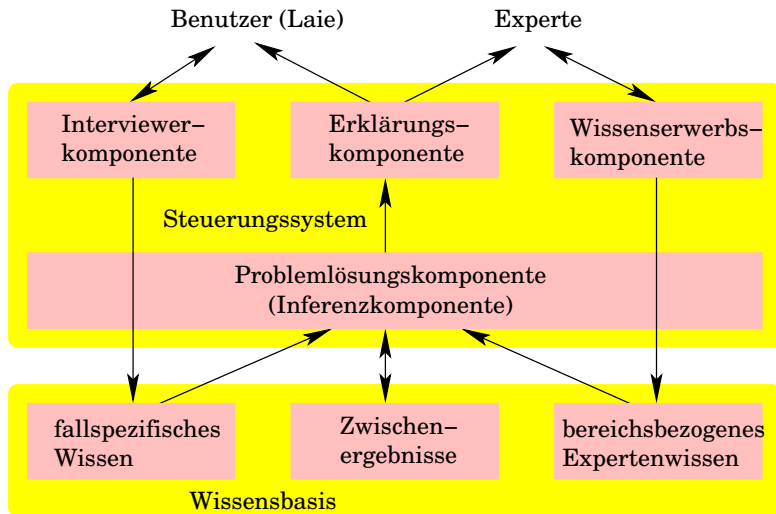
Im Mittelalter wusste man, dass die Erde eine Scheibe ist!

Im weiteren Sinn ist jede Programmierung **Wissensverarbeitung**:

- Jedes Programm kodiert und verarbeitet Wissen.

Im engeren Sinn wird die Wissensverarbeitung der Künstlichen Intelligenz zugeordnet.

Trennung von Wissen und Verarbeitung:



Steuerungssystem:

- *Problemlösungskomponente (Inferenzkomponente):* Interpretiert das Expertenwissen, um eine Lösung des vom Benutzer gegebenen Problems zu finden.
- *Interviewerkomponente:* Führt den Benutzerdialog, liest die Messdaten ein und stellt fallspezifisches Wissen dar.
- *Erklärungskomponente:* Liefert dem Benutzer eine Begründung für die Problemlösung. Manche Menschen akzeptieren die von einer Maschine vorgeschlagene Lösung nur dann, wenn das Programm erklärt, wie es zu der Lösung gekommen ist.
- *Wissenserwerbskomponente:* Aufbau und Wartung des Fachwissens.

Wissensbasis:

- *Expertenwissen*: Fachwissen aus dem Anwendungsbereich, häufig in Form von Regeln.
- *fallspezifisches Wissen*: Daten, die der Anwender zu seinem konkreten Problem beisteuert, häufig in Form von Fakten.
- *Zwischenergebnisse / Problemlösungen*: Werden von der Inferenzkomponente erzeugt.

Vorteil dieser Architektur:

- Die Wissensbasis ist austauschbar!

Arten von Wissen:

- Relationales Wissen spielt im Alltag eine große Rolle:
 - Person X ist mit Y verheiratet.
 - Der Motor ist ein Teil vom Auto.
- Vererbung von Eigenschaften:
 - Der Golf ist ein Auto. Ein Auto hat einen Motor.
 - ⇒ Der Golf hat einen Motor.
- Prozedurales Wissen schreibt eine Folge von Aktionen vor:
 - Eröffnungen beim Schach
 - das Betanken eines Autos
 - das Betreten eines Restaurants
- Logisches Wissen:
 - Jeder Mensch ist sterblich. Sokrates ist ein Mensch.
 - ⇒ Sokrates ist sterblich.

- Zeitabhängiges Wissen

- „Es ist kalt“ ist nicht immer richtig.
 - ... täglich außer Montags ...
 - Devisenkurse, Füllstände usw. sind zeitabhängig.
 - Beim Planen sind zeitliche Relationen wichtig: *vor*, *nach*, ...
- Temporallogik, Situationskalkül

- Unvollständiges Wissen

- Eine Kamera erfasst nicht die gesamte Umgebung. Was liegt hinter dem nächsten Hügel?
 - Nicht alle Zusammenhänge in der Medizin sind bekannt.
- Fehlende Informationen werden durch Annahmen ersetzt.

- Vages Wissen

- unscharf: Der Mann ist groß. Das Wetter ist gut.
 - unsicher: Temperatursensor hat Toleranz von $\pm 1^\circ\text{C}$.
- Fuzzy-Logik, Sicherheitsfaktoren

Darstellung von Wissen:

- Neuronale Netze
- Regeln
- Semantische Netze und Frames
- Aussagen- und Prädikatenlogik
- Skripte
- Bayessche Netze

Anforderungen an die Wissensrepräsentation:

- Wissen ist leicht überprüfbar und aktualisierbar.
- Ableiten neuen Wissens / Schlussfolgern ist möglich.
- Hypothesen können bestätigt oder widerlegt werden.

Warum ist die Darstellung von Wissen wichtig?

- Wissen muss effizient verarbeitet werden.
- Qualität der Antwort muss den Erfordernissen entsprechen.
- Beispiel: Rechnen mit römischen Zahlen vs. indischen Zahlen.

Closed-World-Assumption:

- Nur positive Fakten werden explizit gespeichert. Fehlende Information wird als Falschheit der Information gedeutet.
- Oft richtig und vernünftig, um Wissensbasis klein zu halten:
 - Wenn kein Direktflug von Bonn nach London in der Datenbank gespeichert ist, dann gibt es auch keinen.
 - Die Kannibalen¹ dürfen nicht unter Drogen gesetzt werden und es gibt auch keine Brücke über den Fluss.

¹Drei Missionare und 3 Kannibalen wollen einen Fluss überqueren. Das Ruderboot für die Überfahrt kann nur 2 Personen aufnehmen. Es dürfen nie mehr Kannibalen als Missionare an einem Ort sein, sonst wird es unschön.

Anwendung bei Problemen, die keine algorithmische Lösung haben, z.B. einige Suchprobleme mit exponentiell großem Lösungsraum.

- Logik-Rätsel wie Sudoku: Fülle das Gitter mit den Ziffern 1 bis 9 so, dass jede Ziffer in jeder Spalte, in jeder Zeile und in jedem Block genau einmal vorkommt.
- Zahlenrätsel: SEND + MORE = MONEY
- Spiele: Tic-Tac-Toe, Mühle, Schach, Go

	3							
			1	9	5			
	9	8					6	
8				6				
4					3			1
				2				
	6					2	8	
			4	1	9			5
							7	

Übung: Schreiben Sie ein Programm, das ein gegebenes $n \times n$ -Sudoku für $n = 9, 16, 25, 36, 49$ mittels Backtracking löst. Welche Laufzeiten ergeben sich für große Werte von n ?

Methode der Wissensbasierten Systeme: Einsatz von Wissen!

- Erfolg ist nicht garantiert, aber oft besser als ohne Wissen.
- Wissen ist anwendungsspezifisch, oft als Heuristik realisiert.

Viele bekannte Spiele sind schwer:

- Sudoku ist NP-vollständig: Yato und Seta (2003)
- Mastermind ist NP-vollständig: Stuckman und Zhang (2006)
- Minesweeper ist NP-vollständig: Kaye (2000)
- Brett-Solitär ist NP-vollständig: Uehara und Iwata (1990)
- Sokoban ist PSPACE-vollständig: Dor und Zwick (1999)
- Rush Hour ist PSPACE-vollständig: Flake und Baum (2002)

Eine schöne Übersicht über die Komplexität von Spielen finden Sie unter:

<http://www.cs.wmich.edu/~elise/courses/cs431/icga2008.pdf>

Reduktion des Suchraums durch Wissen:

Streiche in jedem Feld aus der Liste der prinzipiell möglichen Zahlen alle die Zahlen heraus, die bereits in der jeweiligen Zeile, der Spalte oder in dem Unterblock vorkommen.

4			2	7		6		
7	9	8	1	5	6	2	3	4
	2		8	4				7
2	3	7	4	6	8	9	5	1
8	4	9	5	3	1	7	2	6
5	6	1	7	9	2	8	4	3
	8	2		1	5	4	7	9
	7			2	4	3		
		4		8	7			2

⇒

4	¹ ₅	³ ₅	2	7	³ ₉	6	¹ _{8 9}	⁵ ₈
7	9	8	1	5	6	2	3	4
¹ ₆	2	³ _{5 6}	8	4	³ ₉	¹ ₅	¹ ₉	7
2	3	7	4	6	8	9	5	1
8	4	9	5	3	1	7	2	6
5	6	1	7	9	2	8	4	3
³ ₆	8	2	³ ₆	1	5	4	7	9
¹ _{6 9}	7	⁵ ₆	⁶ ₉	2	4	3	¹ _{8 6}	⁵ ₈
¹ _{3 6 9}	¹ ₅	4	³ _{6 9}	8	7	¹ ₅	¹ ₆	2

Naked Single - einzig möglicher Wert: Das Kästchen muss diesen Wert annehmen, da alle anderen Werte bereits vergeben sind.

1			2					4
	9	8	5		3	6	7	
	4			7				
	1						5	9
		2				3		
9	7						2	
				1			6	
	8	5	3		9		1	
4					6			5

Hidden Single - einzig möglicher Ort: Das Kästchen muss diesen Wert annehmen, da alle anderen Kästchen diesen Wert nicht annehmen können.

						5		
1	6		9					
		9		6	4			
								4
4				2		1		
			3	4			5	
		2		8	9			
	1		2	5			3	
7			1					9

Naked Pair: Kommen in zwei Zellen eines Bereichs ausschließlich die zwei gleichen Werte vor, dann müssen sich diese zwei Zellen diese zwei Werte teilen.

4			2	7		6		
7	9	8	1	5	6	2	3	4
	2		8	4				7
2	3	7	4	6	8	9	5	1
8	4	9	5	3	1	7	2	6
5	6	1	7	9	2	8	4	3
³ ₆	8	2	³ ₆	1	5	4	7	9
¹ ₆	7	⁵ ₆	⁶ ₉	2	4	3	¹ ₆	⁵ ₈
¹ ₃	¹ ₅	4	³ ₆	8	7	¹ ₅	¹ ₆	2

Hidden Pair: Kommt ein Wertepaar nur in zwei Zellen vor, dann müssen sich diese beiden Zellen diese Werte teilen, da keine andere Zelle sie übernehmen kann.

	6		3	9		1		
		3	1	5			9	
1	9		4	2	6	3		
8	3		5	7	9	4	1	² ₆
9				6	1	² ₇	³ ₅	² ₃
	5	1		4	3	² ₇	⁶ ₈	9
4	1	9	6	3	5	8	2	7
	2		9	8	4	5		1
	8		7	1	2	9	4	

Weitere Methoden zur Einschränkung des Suchraums bei Sudoku finden Sie bspw. unter:

www.sudokuoftheday.com/pages/techniques-overview.php

Übung: Bauen Sie obige Techniken in Ihren Sudoku-Solver ein und vergleichen Sie die Laufzeiten der beiden Versionen.

Zahlenrätsel: Wissen über Addition beschränkt den Suchraum.

- Bei Addition zweier Ziffern ist der Übertrag 1.

$$\begin{array}{rcccccc} & S & E & N & D & \rightarrow & S & E & N & D \\ + & M & O & R & E & & + & 1 & O & R & E \\ \hline M & O & N & E & Y & & 1 & O & N & E & Y \end{array}$$

Zahlenrätsel: Wissen über Addition beschränkt den Suchraum.

- Bei Addition zweier Ziffern ist der Übertrag 1.

$$\begin{array}{rcccccc} & S & E & N & D & \rightarrow & & S & E & N & D \\ + & M & O & R & E & & + & 1 & O & R & E \\ \hline M & O & N & E & Y & & 1 & O & N & E & Y \end{array}$$

- S muss gleich 8 oder 9 sein, da sonst kein Übertrag erzielt würde. Wir setzen zunächst S auf 9.

$$\begin{array}{rcccccc} & S & E & N & D & \rightarrow & 9 & E & N & D \\ + & 1 & O & R & E & & + & 1 & O & R & E \\ \hline 1 & O & N & E & Y & & 1 & O & N & E & Y \end{array}$$

Zahlenrätsel: Wissen über Addition beschränkt den Suchraum.

- Bei Addition zweier Ziffern ist der Übertrag 1.

$$\begin{array}{rcccccc} & S & E & N & D & \rightarrow & & S & E & N & D \\ + & M & O & R & E & & + & 1 & O & R & E \\ \hline M & O & N & E & Y & & 1 & O & N & E & Y \end{array}$$

- S muss gleich 8 oder 9 sein, da sonst kein Übertrag erzielt würde. Wir setzen zunächst S auf 9.

$$\begin{array}{rcccccc} & S & E & N & D & \rightarrow & & 9 & E & N & D \\ + & 1 & O & R & E & & + & 1 & O & R & E \\ \hline 1 & O & N & E & Y & & 1 & O & N & E & Y \end{array}$$

- Dann muss O gleich 0 sein. (1 ist bereits durch M belegt!)

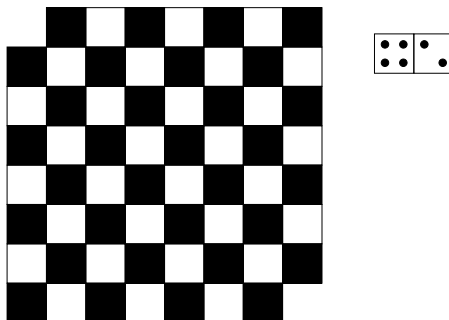
$$\begin{array}{rcccccc} & 9 & E & N & D & \rightarrow & & 9 & E & N & D \\ + & 1 & O & R & E & & + & 1 & 0 & R & E \\ \hline 1 & O & N & E & Y & & 1 & 0 & N & E & Y \end{array}$$

Zahlenrätsel: Was bringt der Einsatz von Wissen?

- Wir haben insgesamt 8 verschiedene Buchstaben im obigen Rätsel.
 - Jeder Buchstabe stellt eine Ziffer dar, wobei verschiedene Buchstaben auch verschiedene Ziffern darstellen.
 - Wir haben also im ursprünglichen Problem $10 \cdot 9 \cdot 8 \cdot 7 \cdot \dots \cdot 3 = 1.814.400$ Kombinationen zu testen.
 - Durch den Einsatz des Wissens legen wir die Werte von M und O eindeutig fest, für S verbleiben nur 2 Möglichkeiten.
- Damit reduzieren wir den Suchraum durch den Einsatz von Wissen auf $2 \cdot 7 \cdot 6 \cdot 5 \cdot \dots \cdot 3 = 5.040$ Kombinationen.

Explizites Wissen über einen Aufgabenbereich verkürzt die Suche!

Plazierung von Dominosteinen auf einem Schachbrett, bei dem zwei diagonal gegenüberliegende Eckfelder fehlen.



Frage: Kann man das Schachbrett so mit Dominosteinen bedecken, dass alle Felder belegt sind, keine Dominosteine über den Rand herausragen und sich keine Dominosteine überlappen?

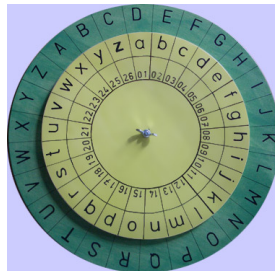
Lösungsmöglichkeiten:

- Vollständige Suche:
Teste alle Möglichkeiten, Dominosteine auf das Brett zu legen. Stop, sobald alle Felder belegt sind oder falls alle Möglichkeiten fehlgeschlagen sind.
- Reduktion des Suchraums durch Symmetriebetrachtungen oder Heuristiken.
- Einsatz von Wissen:
Das Brett hat 30 weiße und 32 schwarze Felder. Jeder Stein bedeckt genau ein schwarzes und ein weißes Feld. $\Rightarrow$ Es gibt keine Lösung für das Problem!

Der Einsatz von Wissen kann eine Suche enorm verkürzen!

Cäsar-Code: Ersetzte jeden Buchstaben des Textes durch den Buchstaben, der im Alphabet n Stellen danach kommt.

Beispiel: Für $n = 2$ wird aus dem Klartext ELEFANTEN der verschlüsselte Text GNGHCPVGP.



Codebrechen mittels Häufigkeitstabelle:

Buchstabe	Häufigkeit
E	17,40%
N	9,78%
I	7,55%
S	7,27%
R	7,00%

Paare	Häufigkeit
ER	4,09%
EN	4,00%
CH	2,42%
DE	1,93%
EI	1,87%

Codierter Text: FKGUKUVGKPGIGJGKOGPCEJTKEJV

Klartext: ??

Buchstabe	Anzahl	relative Häufigkeit
G	6	23,1%
K	5	19,2%
J	3	11,5%
C	2	7,7%
U	2	7,7%

Vermutlich $G \rightarrow E$: DIESISTEINEGEHEIMENACHRICHT

Was ist künstliche Intelligenz?

- 1955 John McCarthy:
Ziel der KI ist es, Maschinen zu entwickeln, die sich verhalten, als verfügten sie über Intelligenz.
- 1991 Encyclopedia Britannica:
KI ist die Fähigkeit digitaler Computer oder computergesteuerter Roboter, Aufgaben zu lösen, die normalerweise mit den höheren intellektuellen Verarbeitungsfähigkeiten von Menschen in Verbindung gebracht werden.
- 1983 Elaine Rich:
Forschung darüber, wie Computer Dinge machen, die Menschen derzeit noch besser beherrschen.

Heute:

- KI ist die Untersuchung von Berechnungsverfahren, die es ermöglichen, wahrzunehmen, zu schlußfolgern und zu handeln.
 - Ermitteln, welche Annahmen über die Repräsentation und Verarbeitung von Wissen und den Aufbau von Systemen die verschiedenen Aspekte der Intelligenz erklären können.
- ⇒ Synthese vor Analyse: Lerne menschliche Intelligenz zu verstehen durch die Konstruktion intelligenter Systeme.

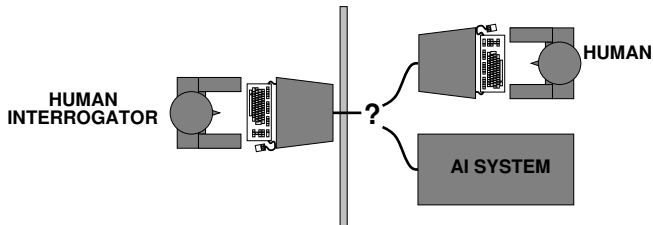
Vorherrschender Gedanke:

- Intelligenz emergiert aus der Interaktion vieler einfacher Prozesse im Zusammenspiel und
- Prozessmodelle intelligenten Verhaltens können mit Hilfe des Computers im Detail untersucht werden.

- Motivation
- *Turing-Test*
- neuronale Netze
- Agenten
- Entscheidungsbäume
- fallbasiertes Schließen

Turing-Test

Der Britische Mathematiker Alan Turing schlug 1950 in seinem Artikel „Computing Machinery and Intelligence“ einen Test zur Überprüfung intelligenten Verhaltens einer Maschine durch Vergleich mit einem Menschen vor:



Der Turing-Test:

- Eine Person befragt einen Menschen und ein KI-System.
- Der Tester kann jede beliebige Frage stellen.
- Der Tester kommuniziert nur indirekt über ein Textmedium.
Dadurch soll vermieden werden, dass der Tester durch das äußere Erscheinen oder durch mechanische Eigenschaften wie eine synthetische Stimme beeinflusst wird.
- Tester soll allein aus den Antworten auf seine Fragen entscheiden, wer Mensch und wer Computer ist.
- Wenn er das nicht kann, ist die Maschine intelligent.

Anforderungen an das Computerprogramm:

- Sprachverstehen und Sprachgenerierung
- Repräsentation von Wissen
- Schlussfolgern
- Lernen

Der Artikel von Alan Turing enthält außerdem erste Gedanken zum Thema Lernen.

Implementieren Sie das Spiel Mastermind. Anstelle der verschieden farbigen Stecker des Originals verwenden wir Ziffern von 0 bis 5.

- Der menschliche Gegenspieler des Computers denkt sich eine Kombination von vier Ziffern aus.
- Das Programm soll mit möglichst wenigen Versuchen diese Kombination ermitteln. In jedem Schritt gibt das Programm eine Kombination aus und wartet auf die Bewertung seines Versuchs durch den menschlichen Gegenspieler.
 - schwarz: Ein Spielstein mit der richtigen Farbe wurde an der richtigen Stelle plaziert.
 - weiß: Ein Spielstein hat zwar die richtige Farbe, steht aber an der falschen Position.
- Durch Analysieren der Bewertungen können nun Informationen gewonnen werden, die helfen, den nächsten Zug zu verbessern.

Beispiel: Der Mensch wählt die Kombination 0,1,2,3. In dem folgenden Dialog macht der Computer jeweils einen Vorschlag, worauf der menschliche Gegenspieler seine Bewertung abgibt.

1. Versuch: 3,3,3,1
--> Bewertung (s,w) ? 0,2
2. Versuch: 5,1,1,3
--> Bewertung (s,w) ? 2,0
3. Versuch: 2,1,4,3
--> Bewertung (s,w) ? 2,1
4. Versuch: 4,1,0,3
--> Bewertung (s,w) ? 2,1

Lösung: 0,1,2,3

Würden Sie ein solches Programm als intelligent bezeichnen?

Joseph Weizenbaum (1966) wird zum Gesellschaftskritiker wegen des Erfolgs seines Programms Eliza:

- Oberflächliche Simulation eines Psychotherapeuten.
- Das Kommunikationsverhalten von Versuchspersonen gegenüber dem Programm entsprach demjenigen gegenüber einem menschlichen Gesprächspartner.
- Die Versuchspersonen waren zum Teil überzeugt, dass der Gesprächspartner ein tatsächliches Verständnis für ihre Probleme aufbrachte.
- Praktizierende Psychiater glaubten ernsthaft daran, mit dem Programm zu einer automatisierten Form der Psychotherapie gelangen zu können.

- Ein Aussagesatz des Benutzers kann in einen Fragesatz umgewandelt werden, indem ein „Warum“ davor gestellt wird.
 - aus „Ich fühle mich unwohl in Gegenwart meiner Mutter.“ wird
 - „Warum fühlen Sie sich unwohl in Gegenwart Ihrer Mutter?“
- Eltern wissen, dass man durch Warum-Fragen schnell an die Wand gespielt wird. Daher: Reagiere auf Schlüsselwörter.
 - auf das Stichwort „Arbeit“ kann das Programm reagieren mit
 - „Haben Sie das Gefühl, Sie arbeiten zuviel?“ oder
 - „Kommt Ihre Familie zu kurz, wenn Sie viel arbeiten?“ oder
 - „Definieren Sie sich über Arbeit? Haben Sie keine Hobbies?“
- Zum Auflockern des Dialogs werden Phrasen eingestreut.
 - „Können Sie darauf bitte genauer eingehen?“
 - „Nehmen Sie kein Blatt vor den Mund.“
 - „Wie fühlen Sie sich jetzt?“
 - „Lassen Sie uns noch etwas bei diesem Thema bleiben.“
 - „Was denken Sie, sind die Gründe dafür?“
 - „Was könnte denn die Ursache von all dem sein?“

Anmerkungen:

- Das Erkennen von Schlüsselwörtern muss auch bei Rechtschreibfehlern korrekt funktionieren. Wir brauchen Fehlertoleranz im Dialog!
- Das Umstellen von Satzteilen ist je nach verwendeter Sprache unterschiedlich kompliziert und erfordert ein Erkennen der Satzstruktur, also der Grammatik.
- Deklination und Konjugation müssen berücksichtigt werden, um Schlüsselwörter zu erkennen: er wird sterben, sie ist gestorben, er starb, du stirbst usw.
- Wir müssen Synonyme erkennen und in die Antworten einbauen: umkommen, versterben, entschlafen, den Tod finden, von uns gehen, ums Leben kommen, dem Schöpfer gegenüber treten usw.

Anmerkungen:

- Um authentisch zu wirken, muss dem Programm die Umgangssprache bekannt sein: abkratzen, draufgehen, verrecken, den Löffel abgeben, über den Jordan gehen usw.
- Die Reaktion auf Schlüsselwörter kann von der Bedeutung eines Satzes abhängen:
 - „Mein *Vater* ist im *Krieg* gefallen“ ist sicherlich anders zu bewerten als „der *Krieg* ist der *Vater* aller Dinge“.
 - „Ich habe das Auto meines Vaters *gestohlen*“ unterscheidet sich deutlich von „ich *stehle* meiner Schwester die Show“.

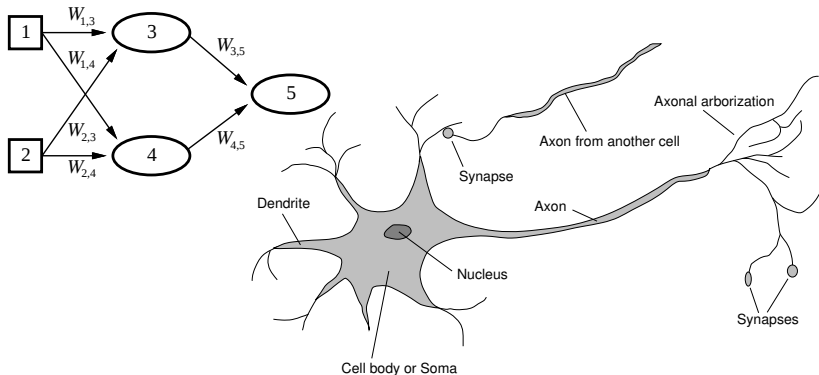
→ Die Semantik ist wichtig und muss erkannt werden!

Um den Turing-Test wirklich zu bestehen, muss schon einiger Aufwand betrieben werden. Wird es jemals gelingen?

- Motivation
- Turing-Test
- *neuronale Netze*
- Agenten
- Entscheidungsbäume
- fallbasiertes Schließen

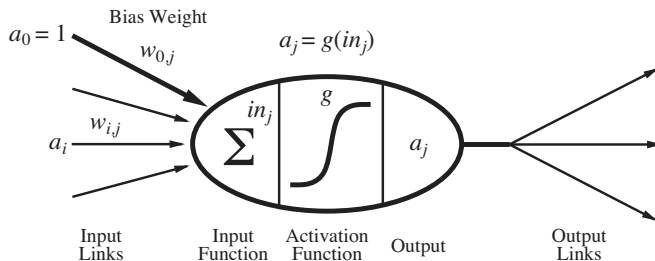
Cullach/Pitt 1943: Integration dreier Quellen zum Vorschlag eines Netzes künstlicher Neuronen: Physiologie der Neuronen, Aussagenlogik und Theorie der Berechenbarkeit (Turing)

Minsky/Edmonds 1951: Build the first artificial neural network that simulated a rat finding its way through a maze.



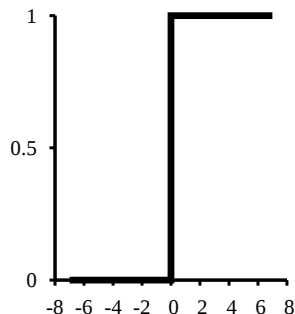
Die von Minsky und Papert (1969) gezeigten theoretischen Grenzen von Perceptrons führen zu einem weitgehenden Ende der Forschungsförderung im Bereich neuronaler Netze.

Künstliches Neuron:

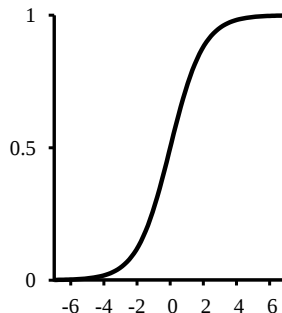


Aktivierungsfunktionen bei neuronalen Netzen

Schwellwertfunktion:

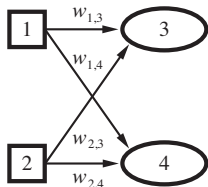


Logistische Funktion:

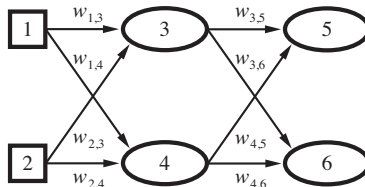


Häufig wird die logistische Funktion $\frac{1}{1+e^{-x}}$ oder der Tangens Hyperbolicus $\tanh(x)$ als Aktivierungsfunktion genutzt.

Vorwärts gerichtete neuronale Netze



(a)

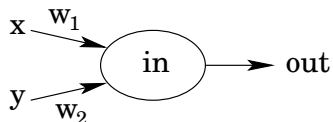


(b)

- **überwachtes Lernen:** Vergleiche Ist- mit Sollausgabe zu einer Eingabe. Schließe aus der Differenz auf die vorzunehmenden Änderungen der Gewichte. → Delta-Regel, Backpropagation
- **unüberwachtes Lernen:** Erfolgt ausschließlich durch Eingabe der zu lernenden Muster. Das Netz verändert sich entsprechend den Eingabemustern von selbst. Oft genutzte Pfade werden verstärkt. → Hebb

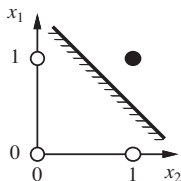
Perceptrons: single layer of artificial neurons connected by weights to a set of inputs

Linear separierbare Funktionen:

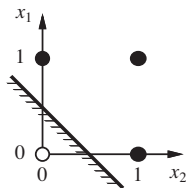


- Ein einzelnes Neuron mit zwei Eingängen x und y liefert $in = x \cdot w_1 + y \cdot w_2$.
- Die Gleichung ist linear in x und y , daher liegen alle Werte, die die Gleichung erfüllen, auf einer Geraden.
- Alle Werte auf einer der Seiten der Geraden produzieren Werte größer als der Schwellwert und aktivieren das Neuron.

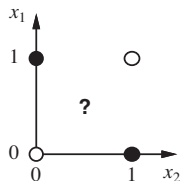
Minsky und Papert: a single layer perceptron cannot simulate a simple exclusive-or function.



(a) x_1 and x_2



(b) x_1 or x_2



(c) x_1 xor x_2

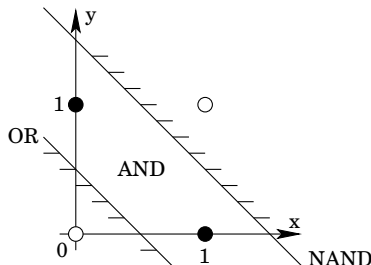
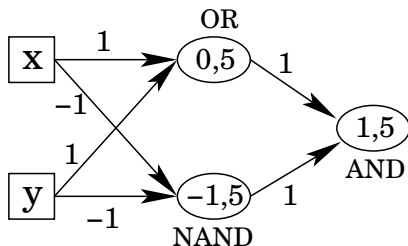
Exklusiv-Oder-Problem:

	w_1	w_2	Schwellwert
AND	1	1	1,5
OR	1	1	0,5
NAND	-1	-1	-1,5
XOR	?	?	?

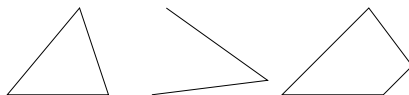
- w_1 und w_2 beeinflussen die Steigung der Geraden,
- der Schwellwert beeinflusst die Position.

Fehler von Minsky und Papert:

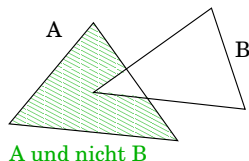
- Sie haben von der Leistung eines einstufigen Perceptrons mit einer trainierbaren Schicht unzulässiger Weise auf die Leistung eines mehrstufigen Perceptrons mit mehreren trainierbaren Schichten geschlossen.
- Mehrstufige Perceptrons sind prinzipiell in der Lage, jede berechenbare Funktion zu realisieren.



Einfache, zweischichtige Netze führen zu konvexen Regionen, die offen oder geschlossen sein können.



Mittels dreischichtigen Netzen können auch nicht-konvexe Regionen dargestellt werden.



Typische Anwendungen: Klassifizierungen von Daten, Prognose für Börsenwerte oder Energieverbrauch, Steuerungen in der Robotik oder von Prothesen, Zeichenerkennung.

Das Buch von Rumelhart und McClelland *Parallel Distributed Processing*, 1986 verbreitet neue Erfolge mit dem wiederentdeckten Backpropagation-Algorithmus.

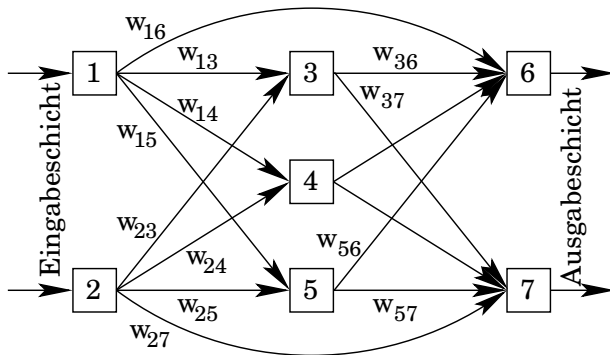
Der Backpropagation-Algorithmus wurde bereits 1969 von Bryson und Ho publiziert. Also im gleichen Jahr, in dem das kritische Buch von Minsky und Papert erschienen ist.

Backpropagation: Gradientenabstiegsverfahren zur Minimierung des Fehlers, also der Abweichung zwischen Ausgabe (Ist-Wert) und erwarteter Ausgabe (Soll-Wert).

- vorwärts gerichtete Netzwerke
- Matrix w speichert Gewichte der Kanten: w_{ij} ist das Gewicht der Kante von Neuron i zu Neuron j
- Aktivierungsfunktion: $\frac{1}{1+e^{-x}}$ ist überall differenzierbar

künstliche neuronale Netze

- out_i : Ausgabe des Neurons i .
- $in_j = \sum_i out_i \cdot w_{ij}$
- die Aktivierungsfunktion liefert: $out_j = \frac{1}{1+e^{-in_j}}$



Oft definiert man den Fehler des neuronalen Netzes

$$E = \frac{1}{2} \sum_j (t_j - out_j)^2$$

als Quadrat der Abweichung zwischen erwarteter Ausgabe und realer Ausgabe, damit sich negative und positive Abweichungen nicht aufheben.

Der Fehler wird minimiert, indem jedes Gewicht um einen durch die Lernrate η festgelegten Bruchteil des negativen Gradienten der Fehlerfunktion geändert wird:

$$\Delta w_{ij} = -\eta \cdot \frac{\partial E}{\partial w_{ij}}$$

Die Lernrate η steuert den Grad der Gewichtsänderung und das negative Vorzeichen bewirkt eine Veränderung entgegen dem Kurvenanstieg in Richtung eines Tals der Fehlerkurve.

Die Kettenregel liefert:

$$\Delta w_{ij} = -\eta \cdot \frac{\partial E}{\partial w_{ij}} = -\eta \cdot \frac{\partial E}{\partial in_j} \cdot \frac{\partial in_j}{\partial w_{ij}}$$

- $\frac{\partial in_j}{\partial w_{ij}} = \frac{\partial}{\partial w_{ij}} \sum_k out_k \cdot w_{kj} = out_i$
- Wir definieren als Fehlersignal: $\delta_j = -\frac{\partial E}{\partial in_j}$

Damit erhalten wir:

$$\Delta w_{ij} = \eta \cdot \delta_j \cdot out_i$$

Gewichtsänderung:

$$w_{ij}^{neu} := w_{ij}^{alt} + \Delta w_{ij}$$

Wir müssen jetzt nur noch das Fehlersignal $\delta_j = -\frac{\partial E}{\partial in_j}$ bestimmen.

Auch hier können wir wieder die Kettenregel anwenden:

$$\delta_j = -\frac{\partial E}{\partial out_j} \frac{\partial out_j}{\partial in_j}$$

Der zweite Faktor entspricht der ersten Ableitung der Aktivierungsfunktion:

$$\begin{aligned} \frac{\partial out_j}{\partial in_j} &= \frac{\partial}{\partial in_j} \frac{1}{1 + e^{-in_j}} = -\frac{-e^{-in_j}}{(1 + e^{-in_j})^2} = \frac{e^{-in_j}}{(1 + e^{-in_j})^2} \\ &= \frac{1}{1 + e^{-in_j}} \cdot \left(\frac{e^{-in_j}}{1 + e^{-in_j}} \right) \\ &= out_j \cdot \left(\frac{1 + e^{-in_j} - 1}{1 + e^{-in_j}} \right) \\ &= out_j \cdot (1 - out_j) \end{aligned}$$

Für die Berechnung des ersten Faktors $-\frac{\partial E}{\partial out_j}$ benötigen wir eine Fallunterscheidung:

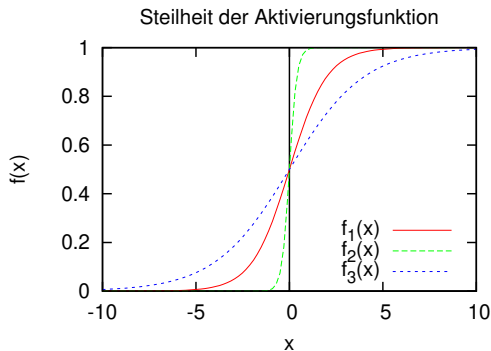
- Für ein **Neuron j der Ausgabeschicht**, dessen Fehler sich nicht auf andere Knoten auswirkt, gilt:

$$-\frac{\partial E}{\partial out_j} = -\frac{\partial}{\partial out_j} \left(\frac{1}{2} \sum_k (t_k - out_k)^2 \right) = (t_j - out_j)$$

- Für ein **Neuron j aus einer Zwischenschicht** ist t_j nicht bekannt. Aber die verallgemeinerte Kettenregel liefert:

$$\begin{aligned} -\frac{\partial E}{\partial out_j} &= -\sum_k \frac{\partial E}{\partial in_k} \frac{\partial in_k}{\partial out_j} \\ &= \sum_k \left(\delta_k \frac{\partial}{\partial out_j} \sum_l out_l \cdot w_{lk} \right) = \sum_k \delta_k w_{jk} \end{aligned}$$

Die Steilheit der Kurve beim Übergang von 0 auf 1 kann durch einen zusätzlichen Faktor im Exponenten gesteuert werden:



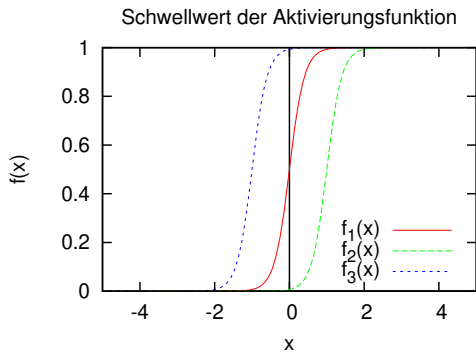
$$f_1(x) = \frac{1}{1 + e^{-x}}$$

$$f_2(x) = \frac{1}{1 + e^{-5x}}$$

$$f_3(x) = \frac{1}{1 + e^{-0.5x}}$$

Durch einen großen Faktor kann also eine Schwellwertfunktion nachgebildet werden.

Soll der Schwellenwert von 0 auf einen anderen Wert verschoben werden, wird noch ein additiver Term im Exponenten eingefügt:



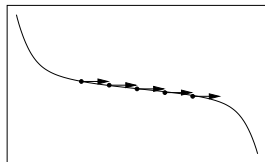
$$f_1(x) = \frac{1}{1 + e^{-5x}}$$

$$f_2(x) = \frac{1}{1 + e^{-5(x-1)}}$$

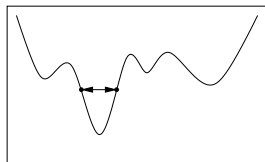
$$f_3(x) = \frac{1}{1 + e^{-5(x+1)}}$$

Festes Eingangssignal 1 über eine Kante mit Gewicht θ mit einem Neuron verbinden. So wird der Schwellenwert zu einem weiteren Gewicht, und wir müssen beim Lernen nur Gewichte optimieren.

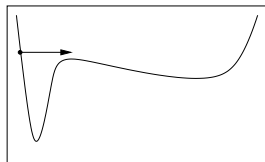
Eventuelle Schwierigkeiten beim Lernen:



Ist die Kurve sehr flach, dann ist der Anstieg betragsmäßig sehr klein, woraus eine geringe Gewichtsänderung resultiert. Der Trainingsprozess nimmt viel Zeit in Anspruch.



Oszillation in sehr steilen Schluchten: Durch den großen Anstieg erfolgt eine große Gewichtsänderung und es wird von einer Wand zur anderen gesprungen.

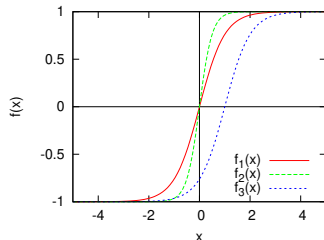


Durch eine steile Wand führt eine große Gewichtsänderung dazu, dass das tiefe Tal übersprungen wird und nur noch ein lokales Optimum gefunden wird.

Lernrate:

- klein: langsame Änderung der Gewichte
 - viele Iterationsschritte notwendig
- groß: schnelle Änderung der Gewichte
 - evtl. Oszillation um idealen Gewichtsvektor
- monoton fallende Folge über alle Lernschritte
 - Problem: Bestimmung der Parameter für diese Folge.

Damit auch inaktive Zustände eine Gewichtsänderung beim Training bewirken, verwendet man den Wertebereich $[-1, +1]$:



$$f_1(x) = \tanh(x)$$

$$f_2(x) = \tanh(2x)$$

$$f_3(x) = \tanh(x - 1)$$

Beispiel: NETtalk² lernte die Aussprache von englischem Text.

- Extrem schwierig für symbolbasierten Ansatz, da die englische Aussprache äußerst unregelmäßig ist.
- vorher: komplizierte Programme liefern schlechte Resultate.
- Aussprache eines Buchstabens ist kontextabhängig, daher wird jeweils ein Textfenster von sieben Zeichen untersucht.
- Während der Text dieses Textfenster durchläuft, gibt NETtalk die Aussprache jeweils des mittleren Buchstabens aus.
- Eigenschaften spiegeln Wesen menschlichen Lernens wider:
 - Je mehr Wörter das Netz aussprechen lernt, desto besser ist es in der Lage, neue Wörter korrekt auszusprechen.
 - Das Netz ist fehlerresistent: Werden einige Gewichte eines gut trainierten Netzes zufällig geändert, bleibt die Aussprache gut.
 - Weniger Neuronen in der verborgenen Schicht als in der Eingabeschicht → irgendeine Art von Abstraktion

²Sejnowski, Rosenberg: Parallel networks that learn to pronounce english text. Complex Systems, 1987.

Counterpropagation:

- Biologisches Vorbild: Man vermutet, dass im Gehirn spezialisierte Module hintereinander geschaltet sind, um eine gewünschte Berechnung auszuführen.
- Hier: Eingabe-, Kohonen- und Grossberg-Schicht.
- vollständige Vernetzung der Schichten: Jedes Neuron der Eingabeschicht ist mit jedem Neuron der Kohonenschicht verbunden. analog: Kohonen- und Grossberg-Schicht.
- normalisierte Eingabevektoren:

$$x_i^{\text{normalisiert}} := \frac{x_i}{\sqrt{x_1^2 + x_2^2 + \dots + x_n^2}}$$

→ Der normalisierte Vektor zeigt in die gleiche Richtung wie der ursprüngliche Vektor, hat aber Einheitslänge.

zugrunde liegende Idee:

- zweidimensionaler Fall: Das Skalarprodukt zwischen dem normalisierten Gewichts- und Eingabevektor ist gleich dem Cosinus des Winkels zwischen Eingabe- und Gewichtsvektor.

$$\vec{a} \cdot \vec{b} = |\vec{a}| \cdot |\vec{b}| \cdot \cos(\angle(\vec{a}, \vec{b}))$$

$$\vec{a} \cdot \vec{b} = a_x \cdot b_x + a_y \cdot b_y$$

Mit Hilfe dieser Formeln können wir den Winkel zwischen den Vektoren bestimmen:

$$\cos(\angle(\vec{a}, \vec{b})) = \frac{a_x \cdot b_x + a_y \cdot b_y}{|\vec{a}| \cdot |\vec{b}|}$$

- Je kleiner der Winkel zwischen Gewichts- und Eingabevektor ist, also je dichter die beiden Vektoren beieinander liegen, umso größer ist der Wert der Aktivierungsfunktion.

zugrunde liegende Idee: (Fortsetzung)

- Das Neuron der Kohonenschicht, dessen Gewichtsvektor die größte Ähnlichkeit zum Eingabevektor zeigt, hat die größte Aktivität.
 - The winner takes it all: Legt man einen Eingabevektor an die Kohonenschicht an, so wird die Ausgabe des Neurons auf 1 gesetzt, dessen Gewichtsvektor die größte Ähnlichkeit mit dem Eingabevektor hat.
- ⇒ Betrachtet man jeden einzelnen Gewichtsvektor als Klasse oder Kategorie, so findet eine Klassifizierung des angelegten Eingabevektors in die für ihn am besten passende Kategorie statt.

Eingabeschicht: Verteilt die Eingangssignale an alle Neuronen der nächsten Schicht, der Kohonen-Schicht.

Kohonen-Schicht:

- Lege einen Eingabevektor an und berechne für alle Neuronen den Wert der Aktivierungsfunktion.
- The winner takes it all: Die Ausgabe des Neurons mit größtem Eingangssignal wird auf 1 gesetzt, alle anderen Ausgaben werden auf 0 gesetzt.
- Die Gewichte des *Gewinner*-Neurons werden korrigiert: Das *Gewinner*-Neuron hat die größte Ähnlichkeit mit dem Eingabevektor; die Gewichtsänderungen machen die Vektoren noch *ähnlicher*.

Training der Kohonen-Schicht: Sei $\vec{w}$ der Gewichtsvektor des *Gewinner*-Neurons, α die Lernrate und $\vec{x}$ der Eingabevektor.

$$\vec{w}^{new} = \vec{w}^{old} + \alpha \cdot (\vec{x} - \vec{w}^{old})$$

Anmerkungen zur Kohonen-Schicht:

- Gewichte werden mit Zufallswerten initialisiert und beim Einsatz der Lernmethode normalisiert.
- Das Kohonen-Netz muss stabil sein, bevor das Grossberg-Netz trainiert werden kann.

Variationen der Kohonen-Schicht:

- „Bewusstseinsparameter“ einführen und in jeder Iteration aktualisieren: Das verhindert, dass einzelne Neuronen zu oft gewinnen. → Alle Knoten tragen zur Repräsentation des Musterraums bei.
- Wähle anstatt eines Gewinners eine Menge von ähnlichsten Knoten und ändere deren Gewichte differenziell.
- Teile den Nachbarknoten des Gewinners eine differenzielle Belohnung zu.

Grossberg-Schicht:

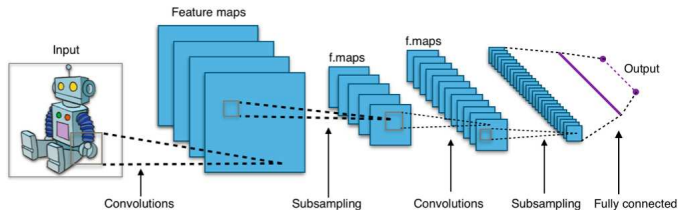
- Die einzige Aufgabe der Grossberg-Schicht ist es, die zuvor gelernten Klassen zu benennen, also die gewünschte Ausgabe zu erzeugen.
- Nach dem Anlegen eines Eingabevektors wird das *Gewinner-Neuron* i der Kohonenschicht bestimmt.
- Anschließend wird das Gewicht v_{ij} zu einem Neuron j der Grossberg-Schicht so geändert, dass der Unterschied zwischen Ist-Ausgabe und Soll-Ausgabe reduziert wird.

Training der Grossberg-Schicht: Sei i das *Gewinner-Neuron*.

$$v_{ij}^{new} = v_{ij}^{old} + \beta \cdot (t_j - v_{ij}^{old})$$

Oft wird $\beta = 0.1$ gesetzt und im Laufe des Trainings, also mit zunehmendem Lernfortschritt, langsam reduziert.

Convolutional Neural Networks sind eine Sonderform von mehrschichtigen Perzeptrons. Daher sind beide Modelle prinzipiell identisch in ihrer Ausdrucksstärke. Bild: wikipedia



Vier Annahmen verringern den Berechnungsaufwand des Netzes und lassen tiefere Netzwerke zu.

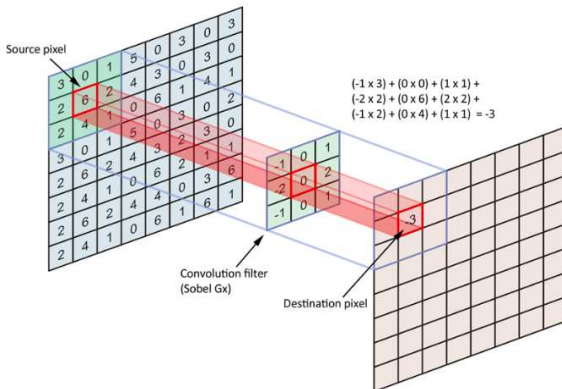
- geteilte Gewichte im Convolutional Layer
- Pooling: verwerfe Großteil der Aktivität eines Layers
- ReLu: extrem einfach zu berechnende Aktivierungsfunktion
- Dropout (entferne zufällige Neuronen) verhindert Overfitting

Convolutional Layer

Der mittels diskreter Faltung ermittelte Input eines jeden Neurons wird von einer Aktivierungsfunktion, üblicherweise Rectified Linear Unit, kurz ReLu $f(x) = \max\{0, x\}$, in den Output verwandelt.

Intuitiv: Bewege schrittweise eine kleine Faltungsmatrix über die Eingabe.

Bild: www.freecodecamp.org

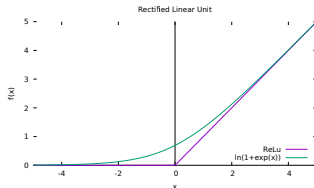


Convolutional Layer

Verschiedene Faltungsmatrizen ergeben unterschiedliche Feature-Maps. Die Gewichte einer Matrix werden mittels Backpropagation trainiert.

Gewichte aller Neuronen eines Convolutional Layers sind identisch
→ schnelleres Training als bei Perzeptrons

Backpropagation verlangt die Berechnung der Gradienten
→ nutze differenzierbare Approximation wie $f(x) = \ln(1 + e^x)$



Berechnung wäre viel zu aufwändig → wähle an der einzigen nicht differenzierbaren Stelle als Steigung Null

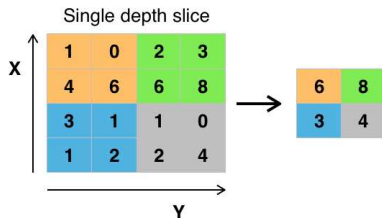
Pooling Layer

Im Pooling-Layer werden überflüssige Informationen verworfen.

Zur Objekterkennung in Bildern etwa, ist die exakte Position einer Kante im Bild von vernachlässigbarem Interesse - die ungefähre Lokalisierung eines Features ist oft hinreichend.

Es gibt verschiedene Arten des Poolings. Mit Abstand am stärksten verbreitet ist das Max-Pooling: Wähle maximale Aktivität aus jedem 2×2 Quadrat.

Bild: wikipedia

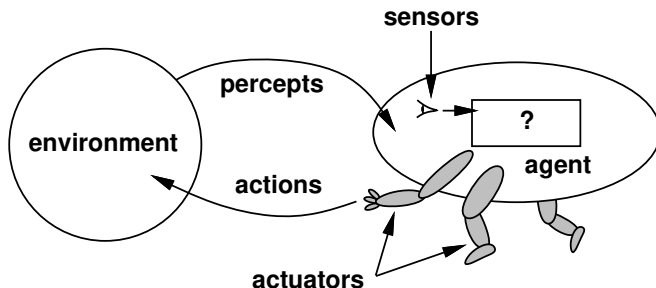


- Motivation
- Turing-Test
- neuronale Netze
- *Agenten*
- Entscheidungsbäume
- fallbasiertes Schließen

Viele Forscher aus den Anwendungsbereichen Robotik, Entwurf von Spielen oder dem Internet haben eine zentrale Wissensbasis oder allgemein gültige Inferenzschemata seit 1995 in Frage gestellt.

Problemlösungsverfahren werden als verteilte Agenten entworfen, die kontextbezogen, autonom und flexibel sind.

Software-Agenten nehmen Umgebung mittels Sensoren wahr und reagieren darauf autonom und rational mittels Effektoren.



Analogie zu agentenbasierten Systemen: Brotbedarf in Städten³

- Es wäre extrem schwierig, ein zentrales Planungsmodul zu entwickeln, dass eine Großstadt erfolgreich mit der vielfältigen Auswahl an täglich frischem Brot versorgen könnte.
- Aber die locker koordinierten Bemühungen der vielen Bäcker, Lastwagenfahrer, Lieferanten der Rohmaterialien sowie der Einzelhändler lösen das Problem sehr gut:
 - Es gibt keinen zentralen Plan.
 - Kein Bäcker hat mehr als ein begrenztes Wissen über den Brotbedarf der Stadt.
 - Jeder Bäcker versucht, sein eigenes Geschäft zu optimieren.
 - Die Lösung des globalen Problems ergibt sich aus den kollektiven Aktivitäten der unabhängigen lokalen Agenten.
- Leider neigen solche Ansätze zur Überproduktion.

³Luger: Künstliche Intelligenz. Pearson Studium.

Software-Agenten:

- Die Problemlösungsaufgabe wird in ihre verschiedenen Bestandteile aufgeschlüsselt.
- Die Teilaufgaben werden kaum oder gar nicht zentral koordiniert.
- Intelligenz ist in einen Kontext eingebettet. → Ein einzelner Problemlöser kann sich einer bestimmten Aufgabe zuwenden, ohne über Informationen zum Status der Lösung in der übergeordneten Problemdomäne verfügen zu müssen.
- Beispiel: Bestellvorgang
 - Ein Agent überprüft die Bestandsinformationen während
 - ein anderer Agent die Kreditwürdigkeit des Kunden prüft.
 - Beide Agenten arbeiten parallel und wissen nichts von der Entscheidung auf höherer Ebene, ob die Bestellung angenommen wird.

Agent:

- Nimmt seine Umgebung über Sensoren wahr.
- Schlussfolgert und entscheidet sich für bestimmte Handlungen.
- Beeinflusst/verändert die Umgebung durch seine Handlungen.

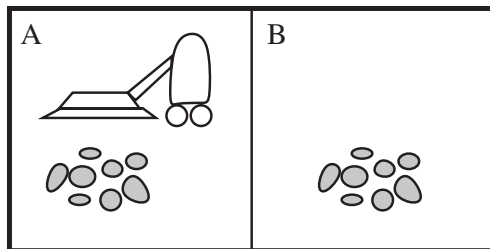
Agenten sind gekennzeichnet durch Rationalität:

- Zur Optimierung des eigenen Nutzens oder um erfolgreich handeln zu können, muss ein Performanzmaß gegeben sein, ein Maß, um den zu erwartenden Erfolg bestimmen zu können.
- Rationalität bedeutet nicht Allwissenheit, rationales Handeln orientiert sich an
 - dem Performanzmaß,
 - aktuellen und früheren Wahrnehmungen,
 - Wissen über die Umgebung,
 - generellen Handlungsmöglichkeiten.

Eigenschaften von Agenten-Systemen:

- Jeder Agent verfügt nur über unvollständige Informationen. Seine Fähigkeiten reichen nicht aus, um das gesamte Problem zu lösen.
 - Es gibt keine globale Instanz zur Steuerung des Systems zur Lösung des Gesamtproblems.
 - Wissen und Eingabedaten liegen in dezentraler Form vor.
 - Schlussfolgerungsverfahren werden asynchron durchgeführt.
- starke Analogie zu objektorientierten Programmen

Staubsauger-Welt⁴:



- Umgebung: Felder A, B, ggf. mit Schmutz
- Sensoren: für Position und Schmutz
- Aktuator: nach rechts/links bewegen; saugen

⁴Russel, Norvig: Artificial Intelligence – A Modern Approach. Prentice Hall

einfaches Agentenprogramm für die Staubsaugerwelt:

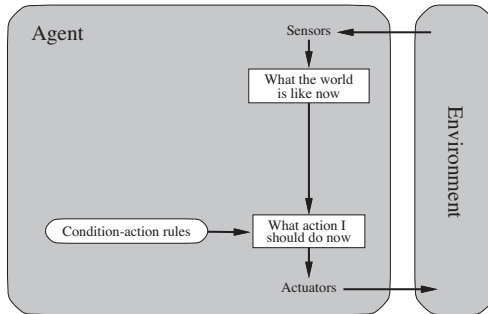
- A oder B schmutzig → saugen
- A sauber → gehe nach rechts
- B sauber → gehe nach links

Anmerkungen:

- Programm berücksichtigt vergangene Wahrnehmungen nicht.
- Wenn alle Felder sauber sind, oszilliert der Sauger zwischen den Feldern.

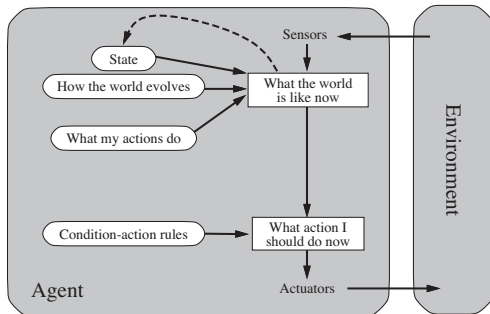
einfache Reflex-Agenten:

- Aktion hängt nur von aktueller Wahrnehmung ab
- Kein Gedächtnis



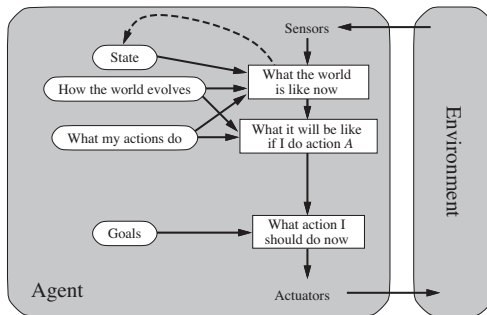
Modellbasierte Agenten:

- Zustandsspeicherung verhindert Oszillation
- Modell der Wirklichkeit beschreibt
 - zeitliche Entwicklung der Umgebung unabhängig vom Agenten
 - die Auswirkungen der Aktionen auf die Welt
- Aber: keine was-wäre-wenn Simulation



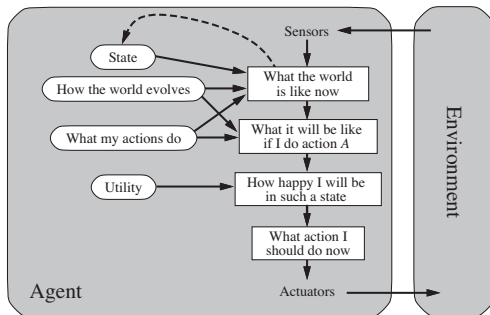
Zielorientierte Agenten:

- Bisher: Ziel wurde implizit durch Regeln kodiert.
- Jetzt: explizite Definition von erstrebenswerten Zielen
- Aktionen nicht aus Regeln ableiten, sondern aus Überlegung „was passiert wenn“ → nutze Planen und Suchstrategien



Nutzenbasierte Agenten:

- Lösungen der zielbasierten Agenten sind oft nicht effizient: Jede Entscheidung, die irgendwie zum Ziel führt, ist ok.
- Nutzenfunktion ordnet aktuellen oder zukünftigen Zuständen Bewertungen über die Nützlichkeit zum Erreichen des Ziels zu.
- Vorteil: Günstige Wege zum Ziel können bevorzugt werden.
- Aber: Aufstellen der Nutzenfunktion verlangt viel Information.



- Betonung der Interoperabilität, auch unter nicht-idealen Bedingungen, z.B. ist Wahrnehmung nie perfekt.
- Erfolge in vielen Teilbereichen verstärken wieder den Blick auf das „Ganze“.

Beispiele:

- Microsofts Heftklammer *Clippy*
- Autonome Staubsauger oder Rasenmäher
- Mobile Informationsagenten

It looks like you're writing a letter.

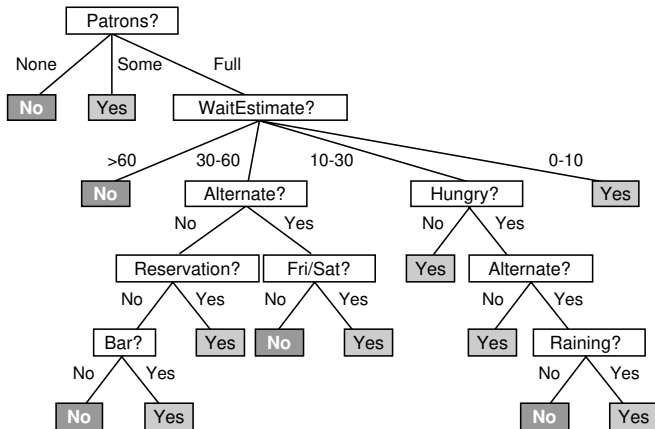
Would you like help?

- Get help with writing the letter
- Just type the letter without help
- ☐ Don't show me this tip again



- Motivation
- Turing-Test
- neuronale Netze
- Agenten
- *Entscheidungsbäume*
- fallbasiertes Schließen

Entscheidungsbäume werden zur automatischen Klassifikation genutzt und damit auch zur Lösung von Entscheidungsproblemen.



```
DECTREE( $E$  : examples,  $A$  : attributes, default : classification)
if  $E = \emptyset$  then
    return default
else if all  $E$  have the same classification  $c$  then
    return  $c$ 
else if  $A = \emptyset$  then
    throw Error           ▷ examples with different classification
else
     $a := \text{CHOOSEATTRIBUTE}(A, E)$ 
     $T :=$  a new decision tree with root  $a$ 
    for all attribute value  $w_i$  of  $A$  do
         $E_i := \{e \in E \mid a(e) = w_i\}$ 
         $T_i := \text{DECTREE}(E_i, A - \{a\}, \text{MAJORITYVAL}(E))$ 
        add a branch to  $T$  with label  $w_i$  and subtree  $T_i$ 
    return  $T$ 
```

Der Aufbau (Induktion) der Bäume erfolgt rekursiv.

Dazu müssen Trainingsdaten vorliegen, also zu jedem Objekt des Datensatzes muss die Klassifikation des Zielattributs bekannt sein.

Bei jedem Induktionsschritt wird das Attribut gesucht, mit welchem sich die Trainingsdaten in diesem Schritt bezüglich des Zielattributs am besten klassifizieren lassen.

Als Maß für die Bestimmung der besten Klassifizierung kommen Entropie-Maße zur Anwendung.

Das ermittelte Attribut wird nun zur Aufteilung der Daten verwendet. Auf die so entstandenen Teilmengen wird die Prozedur rekursiv angewendet, bis in jeder Teilmenge nur noch Objekte mit einer Klassifikation enthalten sind (Greedy Algorithmus).

Generelles Prinzip des Lernens: *Occams Razor*

Bevorzuge die einfachste Hypothese, die konsistent mit allen Beobachtungen ist.

Bei Entscheidungsbäumen wird immer der kleinste Baum gesucht:

Die Chance, dass irgendein kleiner Baum, der insgesamt viele falsche Entscheidungen trifft, zufälligerweise mit allen Beispielen konsistent ist, ist sehr gering. Ein kleiner Baum, der konsistent mit allen Beispielen ist, ist daher eher korrekt als ein großer Baum.

ID3-Algorithmus: Sei S die Menge der Trainingsbeispiele mit ihrer jeweiligen Klassifizierung, sei $a \in A$ das zu prüfende Attribut, sei $V(a)$ die Menge der möglichen Attributwerte von a und sei S_v die Untermenge von S , für die das Attribut a den Wert v annimmt.

Der Informationsgewinn $Gain(S, a)$, der durch Auswahl des Attributs a erzielt wird, errechnet sich dann als Differenz der Entropie von S und der erwarteten Entropie von S bei Fixierung von a :

$$Gain(S, a) = Entropie(S) - \sum_{v \in V(a)} \frac{|S_v|}{|S|} \cdot Entropie(S_v)$$

Wähle Attribut mit dem maximalen Gewinn $\max_{a \in A} \{Gain(S, a)\}$ aus.

Diese Wahl führt zur Bevorzugung von Attributen mit vielen Wahlmöglichkeiten und damit zu einem breiten Baum.

Der Informationsgehalt $I(c)$ eines Zeichens c hängt von der Wahrscheinlichkeit p_c seines Auftretens ab:

$$I(c) = \log_2 \left(\frac{1}{p_c} \right) = \log_2 (p_c^{-1}) = -\log_2(p_c)$$

Shannon definierte die Entropie H als Erwartungswert des Informationsgehalts:

$$\text{Entropie}(S) = H(S) = \sum_{c \in S} p_c \cdot I(c) = - \sum_{c \in S} p_c \log_2(p_c).$$

Setze $0 \cdot \log_2(0) = 0$ entsprechend dem Grenzwert $\lim_{x \rightarrow 0} x \log_2(x)$.

- Motivation
- Turing-Test
- neuronale Netze
- Agenten
- Entscheidungsbäume
- *fallbasiertes Schließen*

Ausgangshypothese: Ähnliche Fälle haben ähnliche Lösungen.

- Immanuel Kant: Alles Wissen stammt aus Erfahrung.
- Konfuzius: Der Mensch hat dreierlei Wege klug zu handeln.
Erstens durch Nachdenken - das ist der edelste, zweitens durch Nachahmen - das ist der leichteste, und drittens durch Erfahrung - das ist der bitterste.

„Käpt'n, Käpt'n, wir haben Wasser im Boot!“ – „Dann zieh' doch den Stöpsel 'raus, damit es ablaufen kann“

Case-Based Reasoning (CBR):

- maschinelles Lernverfahren zur Problemlösung durch Analogieschluss. → kein regelbasierter, sondern erinnerungsbasierter Prozess
- Zentrales Element: Fallbasis, in der bereits gelöste Probleme als Fall gespeichert sind. Ein Fall besteht mindestens aus einer Problembeschreibung und einer zugehörigen Problemlösung.
- Ziehe zur Lösung eines gegebenen Problems die Lösung eines ähnlichen und früher bereits gelösten Problems heran.

Vorteile:

- Experten müssen keine Regeln explizit festlegen sondern Wissen kann implizit dargestellt werden.
- Wissensbasis wird durch Bearbeitung und Speicherung neuer Fälle ständig erweitert und das System lernt automatisch.

Anwendungsbereiche:

- Rechtsprechung: Argumentation erfolgt häufig und explizit nach Sachlage früherer, vergleichbarer Fälle.
- Medizin: Jeder Patient ist ein Fall, auf den später als Erfahrungswissen zurückgegriffen werden kann.
- e-commerce: Verkaufsagenten suchen Produkte aus dem Angebot heraus, indem sie ähnliche Kunden vergleichen.
„Kunden, die Produkt A kauften, kauften häufig auch Produkt B“

Wichtig: Erklärungskomponente ist beim CBR sehr ausgeprägt, Problemlösungen sind sehr gut verständlich.

Qualität eines CBR-Systems ist abhängig von

- seiner Erfahrung, also dem Umfang der Fallbasis
- seiner Fähigkeit, neue Situationen mit gespeicherten Fällen in Verbindung zu bringen
- seiner Fähigkeit, alte Lösungen an neue Probleme anzupassen
- der Art und Weise, wie neue Lösungen evaluiert und bei Bedarf korrigiert werden
- der Integration neuer Erfahrungen in die Fallbasis

Prozess mit vier Phasen:

- *Retrieve*: Ausgehend zur gegebenen Problembeschreibung in der Fallbasis ein möglichst ähnliches Problem ermitteln.
→ Ähnlichkeit der Problembeschreibungen bestimmen
- *Reuse*: Die Lösung des ähnlichsten Falls als ersten Lösungsvorschlag übernehmen.
- *Revise*: Falls das aktuelle Problem nicht genau so zu lösen ist wie das frühere, überprüfe die zuvor gewonnene Lösung und passe sie gegebenenfalls an die konkreten Bedingungen an.
- *Retain*: Der überarbeitete Fall wird in der Fallbasis gespeichert und steht damit für zukünftige Anfragen zur Verfügung.
Auf diese Weise lernt das System mit jedem weiteren gelösten Problem hinzu und verbessert so seine Leistungsfähigkeit.

Das kann algorithmisch nur realisiert werden, wenn wir einen Ähnlichkeitsbegriff definieren.

Repräsentation eines Falls x durch ein Tupel $x = (x_1, \dots, x_n)$, wobei jedes x_i aus dem Wertebereich des i -ten Merkmals stammt.

Bestimme die Gesamt-Ähnlichkeit zweier Fälle $x = (x_1, \dots, x_n)$ und $y = (y_1, \dots, y_n)$ durch einen Abgleich der einzelnen Merkmale:

$$\text{sim}(x, y) = f(\text{sim}_1(x_1, y_1), \dots, \text{sim}_n(x_n, y_n)),$$

Dabei ist sim_i eine reelle Funktion, die die Ähnlichkeit zwischen verschiedenen Werten eines Merkmals bestimmt.

Bei zwei-wertigen Attributen $x_i \in \{0, 1\}$ wird oft die Hamming-Ähnlichkeit genutzt.

$$d_H(x, y) = \sum_{i=1}^n |x_i - y_i| = \sum_{i: x_i \neq y_i} 1$$

Umrechnung in ein Ähnlichkeitsmaß zwischen 0 und 1:

$$\begin{aligned} \text{sim}_H(x, y) &= 1 - \frac{d_H(x, y)}{n} = 1 - \frac{\sum_{i=1}^n |x_i - y_i|}{n} \\ &= \frac{n - \sum_{i=1}^n |x_i - y_i|}{n} = \frac{\sum_{i=1}^n (1 - |x_i - y_i|)}{n} \end{aligned}$$

Für ein Ähnlichkeitsmaß s gilt:

$$s(i, j) = s(j, i) \quad 1 = s(i, i) \geq s(i, j) \quad 0 \leq s(i, j) \leq 1$$

Falls einige Attribute wichtiger als andere sind, verwendet man stattdessen oft die gewichtete Hamming-Ähnlichkeit:

$$\begin{aligned} \text{sim}_H^w(x, y) &= \frac{\sum_{i=1}^n w_i \cdot (1 - |x_i - y_i|)}{\sum_{i=1}^n w_i} \\ &= \frac{\sum_{i=1}^n w_i - \sum_{i=1}^n w_i \cdot |x_i - y_i|}{\sum_{i=1}^n w_i} \\ &= 1 - \frac{\sum_{i=1}^n w_i \cdot |x_i - y_i|}{\sum_{i=1}^n w_i} = 1 - \frac{\sum_{i: x_i \neq y_i} w_i}{\sum_{i=1}^n w_i} \end{aligned}$$

Verallgemeinerte Ähnlichkeiten:

$$\text{sim}(x, y) = \frac{\sum_{i=1}^n w_i \cdot \text{sim}_i(x_i, y_i)}{\sum_{i=1}^n w_i}$$

Repräsentation von Fällen:

- Ein Fall soll Wissen in einem bestimmten Kontext repräsentieren.
- Ein Fall kann Erfahrungen dokumentieren, die von der Norm abweichen oder die einen anderen Verlauf nehmen als erwartet.
- Ein Fall soll dabei helfen, ein gewisses Ziel zu erreichen.
- Ein Fall kann ein warnendes Beispiel liefern.

Ein Fall repräsentiert Wissen im Kontext und dokumentiert eine Erfahrung, die eine wichtige Aussage bzw. Lehre im Hinblick auf gewisse Ziele beinhaltet.

Repräsentation von Wissen wie in KI üblich: Aussagen- und Prädikatenlogik, Regeln, semantische Netze, usw.

Indizierung:

- Indizes enthalten alle wichtigen Angaben, die Fälle voneinander unterscheiden
- Suchalgorithmen nutzen Indizes, um Fälle auszuwählen
- Problem: sicherstellen, dass relevanter Fall auch gefunden wird

Gute Indizes sollen

- hinreichend abstrakt sein, um zu einer Vielzahl zukünftiger Situationen zu passen
 - hinreichend konkret sein, um eine leichte Identifizierbarkeit in zukünftigen Situationen zu gewährleisten
- sorgfältige Situationsanalysen nötig, oberflächliche Faktensammlungen ungeeignet für Indizierung

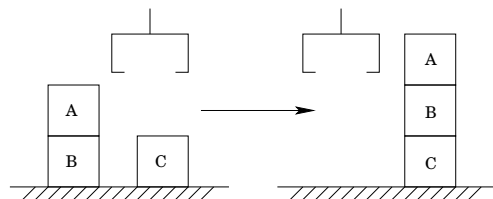
bla bla bla ...

- KI ist etabliert:
 - Die meisten Informatik-Abteilungen haben KI-Professur.
 - Wissensbasierte Techniken werden in der Industrie breit genutzt.
- Relative Selbständigkeit anwendungsorientierter Teilbereiche:
 - Wissenssysteme (Diagnostik, Konfigurierung; Scheduling, Planung, Simulation; Wissensportale)
 - Sprachverarbeitung (hohe Qualität bei gesprochener Sprache; einfaches Verstehen geschriebener Texte - shallowparsing)
 - Bildverarbeitung
 - Robotik (Robo-Cup, Aibo, Marsroboter)
- Betonung logischer Grundlagen:
 - Hidden Markov Models
 - Bayessche Netze
 - logikbasierte Wissensrepräsentationen

- Schachprogramme haben Weltmeisterniveau erreicht.
- Natural Language Processing (NLP):
 - natürlichsprachliche Auskunftssysteme wie z.B. Reisebuchung
 - Spracheingabe, automatische Übersetzer
- erfolgreiche Expertensysteme in vielen Bereichen
 - Selbstdiagnose bei großen Maschinen
 - XUMA: eXpertensystem Umweltgefährlichkeit von Altlasten
Das System soll Fachleute in Behörden und Ingenieurbüros bei der Bewertung von Altlasten und bei der Beurteilung ihrer Umweltgefährdung unterstützen.
 - PROPLANT: Auf Grundlage regionaler Witterung erfolgt eine schlagbezogene Pflanzenschutzberatung. Mittlerweile auch als internetbasierte Version verfügbar.
- automatische Planung und Scheduling in der Raumfahrt
- automatischer Autofahrer, Rasenmäher, Staubsauger
- Dokumentenanalyse: Belegleser
- Fuzzy Control: Fotoapparate, Waschmaschinen, Bremsanlagen

Ziel ist das Aufstellen einer Sequenz von Aktionen, die ausgehend von einem Startzustand einen gewünschten Zielzustand erreichen, unter Berücksichtigung aller Laufzeitbedingungen.

Beispiel: Blocks World



Plan:

```
unstack(A,B)
putdown(A)
pickup(B)
stack(B,C)
pickup(A)
stack(A,B)
```

Zustandsbeschreibung mittels Prädikaten:

- vorher: `onTable(B)`, `onTable(C)`, `upon(A,B)`, `empty()`, `clear(A)`, `clear(C)`
- nachher: `onTable(C)`, `upon(B,C)`, `upon(A,B)`, `empty()`, `clear(A)`

Regeln zur Erzeugung neuer Zustände:

- **pickup(X)**
 - PRECOND: $\text{empty}() \wedge \text{clear}(X)$
 - EFFECT: $\text{hold}(X)$
- **putdown(X)**
 - PRECOND: $\text{hold}(X)$
 - EFFECT: $\text{empty}() \wedge \text{onTable}(X) \wedge \text{clear}(X)$
- **stack(X,Y)**
 - PRECOND: $\text{hold}(X) \wedge \text{clear}(Y)$
 - EFFECT: $\text{upon}(X,Y) \wedge \text{empty}() \wedge \text{clear}(X)$
- **unstack(X,Y)**
 - PRECOND: $\text{upon}(X,Y) \wedge \text{clear}(X) \wedge \text{empty}()$
 - EFFECT: $\text{clear}(Y) \wedge \text{hold}(X)$

Was ist Planen?

- *Wissensbasis*: Definition von Operatoren.
- *Input*: Beschreibung von Start- und Zielsituation.
- *Output*: Der Plan, eine Sequenz von Aktionen.

Anwendungsbeispiele:

- Produktionsplanung
- Reiseplanung
- Spiele: Solitär, Rush Hour, 8-Puzzle usw.

Erzeugen eines Plans:

- Durchsuche einen Raum möglicher Aktionen so lange, bis eine Abfolge gefunden wird, die notwendig ist, um die betreffende Aufgabe zu bewältigen.
- Der Raum repräsentiert Zustände der Welt, die durch anwenden der Aktionen geändert werden.
- Planung verlässt sich auf Suchtechniken.
- Heuristiken können den Suchraum begrenzen.

Probleme:

- Die Beschreibung der Zustände der Welt ist sehr komplex.
- Wie können Pläne verallgemeinert werden?
- Systemwiederherstellung nach unerwarteten Planfehlschlägen?

- Im Laufe der Zeit ändern die Aktionen `unstack`, `stack`, `pickup` und `putdown` die Welt.
- Der Wahrheitswert von Aussagen wie `onTable(A)` ändert sich und ist abhängig von der aktuellen Situation.
- Prädikate, die Merkmale des Problembereichs beschreiben, erhalten als letztes Argument die Situation, auf die sich die Aussage bezieht.
 - Die Konstante `S0` bezeichnet die initiale Situation.
 - Es gibt ein spezielles Funktionssymbol `do`, das aus einer Aktion `a` und einer Situation `s` eine neue Situation `do(a,s)` bildet.
 - `do(stack(B,A),do(stack(A,C),do(unstack(B,C),S0)))` bezeichnet die Situation, die aus der initialen Situation durch Anwenden der Aktionsfolge `[unstack(B,C), stack(A,C), stack(B,A)]` entsteht.

Man nennt eine Situation auch Historie.

⁵Beierle, Kern-Isberner: Methoden wissensbasierter Systeme. Vieweg, 2003.

- Ausführungsbedingungen: Die Formel $\text{Poss}(a, s)$ beschreibt, ob in Situation s die Aktion a ausgeführt werden kann.
 - $\text{Poss}(\text{stack}(x, y), s) \equiv \text{hold}(x, s) \wedge \text{clear}(y, s) \wedge x \neq y$
 - $\text{Poss}(\text{pickup}(x), s) \equiv \text{empty}(s) \wedge \text{clear}(x, s)$
- Effextaxiome spezifizieren die Veränderungen, die eine Aktion bewirkt.
 - $\text{Poss}(\text{stack}(x, y), s) \Rightarrow \text{upon}(x, y, \text{do}(\text{stack}(x, y), s))$
 - $\text{Poss}(\text{stack}(x, y), s) \Rightarrow \neg \text{clear}(y, \text{do}(\text{stack}(x, y), s))$
- Problem: In unserem Beispiel können wir nach der Ausführung von $\text{unstack}(A, B)$ in der initialen Situation nicht auf den Wahrheitswert von $\text{onTable}(B, \text{do}(\text{unstack}(A, B), S_0))$ schließen!

Die Operatoren beschreiben nur, was sich ändert, aber nicht das, was gleich bleibt! $\rightarrow$ Frame-Problem⁶

⁶analogy: Frames of a movie, in which normally very little changes from one frame to the next.

Rahmenproblem:

- Schwierigkeit, die Gültigkeit von Formeln auf die nachfolgende Situation zu übertragen.
- Lösung: Rahmenaxiome beschreiben genau, was durch eine Operation *nicht* verändert wird.
- Beispiel: $\text{pickup}(x)$
 - Zwei Blöcke, die durch die pickup-Operation nicht bewegt wurden, sind immer noch aufeinander:
$$\text{upon}(x,y,s) \wedge x \neq u \Rightarrow \text{upon}(x,y,\text{do}(\text{pickup}(u),s))$$
 - Ein Block befindet sich immer noch auf dem Tisch, wenn er durch die pickup-Operation nicht bewegt wurde:
$$\text{onTable}(x,s) \wedge x \neq u \Rightarrow \text{onTable}(x,\text{do}(\text{pickup}(u),s))$$
- Die Formulierung in Prädikatenlogik ist in der Praxis zu ineffizient: Sehr viele Rahmenaxiome beschreiben sehr viele (scheinbar redundante) Fakten, die erhalten bleiben.

Stanford Research Institute Problem Solver:

- Anfang der 1970er Jahre entwickelt, um die Probleme des Situationskalküls zu vermeiden, ohne jedoch die Logik als Basis von Zustandsrepräsentationen aufzugeben.
- Beeinflusste viele andere Planungssysteme, die in der Folgezeit entstanden.
- Beschreibe Zustände durch Mengen von Formeln, die wir Datenbasis nennen.
- Unique-Names-Assumption: Verschiedene Namen repräsentieren verschiedene Objekte.
- Closed-World-Assumption: Es werden nur positive Fakten gespeichert und alles, was nicht in der Datenbasis vorhanden ist, wird als falsch angesehen.

⁷Beierle, Kern-Isberner: Methoden wissensbasierter Systeme. Vieweg, 2003.

Operatoren wie `unstack(X,Y)` haben folgende Komponenten:

- Menge von Variablen: allquantifiziert
- Vorbedingungsteil: Konjunktion von positiven und negativen Literalen
 - PRECOND: `upon(X,Y)  $\wedge$  clear(X)  $\wedge$  empty()`
- Effekte des Operators: positive und negative Fakten, die die entsprechenden Fakten in der Situation überschreiben.
 - ADD: `clear(Y)  $\wedge$  hold(X)`
 - DELETE: `upon(X,Y)  $\wedge$  clear(X)  $\wedge$  empty()`
- Die Ausführung eines Operators besteht aus zwei Schritten:
 - lösche alle Literale der DELETE-Liste aus der Datenbasis
 - füge alle Literale der ADD-Liste zu der Datenbasis hinzu
 - alle anderen Literale der Datenbasis bleiben erhalten
- Eine konkrete Anwendung eines Operators muss voll instanziiert sein: Alle Variablen sind durch Konstanten substituiert.

Vorwärtssuche: (faktenbasiert)

- Gehe von der initialen Datenbasis aus und
- ändere die Datenbasis durch Operatoranwendungen,
- bis die Datenbasis die Zielbeschreibung erfüllt.

→ Ist in realistischen Anwendungsszenarien nicht praktikabel!

Rückwärtssuche: (zielorientiert)

- Gehe von dem zu erreichenden Ziel aus.
- Um einen Operator anwenden zu dürfen, müssen die Vorbedingungen erfüllt sein.
- Stelle die Vorbedingungen als Teilziel auf.

Mischung aus beiden Verfahren ist wünschenswert, um Suchraum zu beschneiden → means-end-analysis

Mittel-Ziel-Analyse beim General Problem Solver⁸:

- Ermittle syntaktische Unterschiede zwischen dem aktuellen Zustand und dem Zielzustand.
- Wähle einen dieser Unterschiede aus und suche den Operator, der diesen Unterschied reduzieren kann. Probleme:
 - Operator kann im aktuellen Zustand nicht angewendet werden.
→ erzeuge Teilziel, für den Operator angewendet werden kann
 - Operator erzeugt nicht den exakten Zielzustand. → Zwischenzustand wird neuer Startzustand für zweites Problem
- Hoffnung: GPS eignet sich als allgemeine, von Wissensdomäne unabhängige Architektur zum intelligenten Problemlösen, da nur die syntaktische Form von Zuständen untersucht wird.
- Leider: Es scheint nicht nur eine Heuristik zu geben, die sich erfolgreich auf alle Problem-domänen anwenden lässt.

⁸Newell, Shaw, Simon: Report on a general problem-solving program.

Proceedings of the International Conference on Information Processing, 1959

Anmerkungen:

- Linearität:

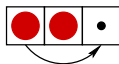
Einem Plan liegt eine totale lineare zeitliche Ordnung zugrunde. Der Planungsalgorithmus nutzt diese Ordnung.

- Lösung des Rahmenproblems:

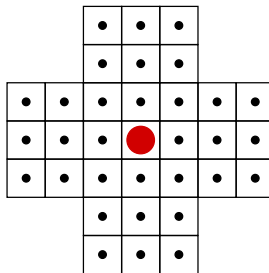
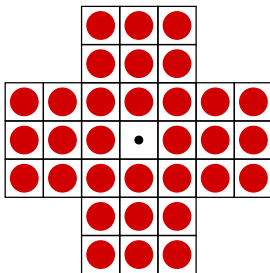
Jeder Operator ändert genau das, was als sein Effekt angegeben ist. Alle nicht angegebenen Fakten bleiben erhalten.

lässt sich nicht in Prädikatenlogik erster Stufe ausdrücken

Solitär: Gegeben ist ein (englisches) Spielbrett in Kreuzform mit 33 Löchern, in denen 32 Stäbchen stecken. Das Loch in der Mitte ist leer. Man muss nacheinander die Stäbchen entfernen, indem man sie in waagerechter oder senkrechter Richtung überspringt.



Am Ende muss ein Stäbchen in der Mitte übrigbleiben.



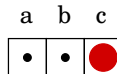
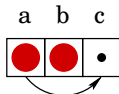
Übung: Schreiben Sie ein Programm, dass zu einem gegebenen Spielbrett eine Zugreihenfolge ermittelt. Welche Laufzeit ergibt sich?

Erkenne aussichtslose Spielsituation mittels Pagoda-Funktion pg :

- Sei P die Menge der zulässigen Positionen beim Solitär. Dann sei $\sigma : P \rightarrow \mathbb{R}$ eine Funktion, so dass für drei benachbarte Positionen a, b, c gilt: $\sigma(a) + \sigma(b) \geq \sigma(c)$
- Sei $B \subset P$ die Menge der aktuell belegten Positionen. Dann ist der Wert der Pagoda-Funktion definiert als

$$pg(B) := \sum_{b \in B} \sigma(b).$$

- Die Werte der Pagoda-Funktion sind monoton fallend im Laufe eines Spiels.



Die Suche kann abgebrochen werden, wenn der Wert der Pagoda-Funktion der aktuellen Spielsituation kleiner ist als der Wert für die Zielsituation.

Übung: Implementieren Sie folgende Pagoda-Funktion und vergleichen Sie die Laufzeiten Ihrer Implementierungen.

```
      -1  0 -1
      1  1  1
-1  1  0  1  0  1 -1
  0  1  1  2  1  1  0
-1  1  0  1  0  1 -1
      1  1  1
      -1  0 -1
```

Beispiel: Produktionsplanung⁹ für zwei Produkte, zu deren Herstellung drei Maschinen A, B und C benötigt werden:

- maximale monatliche Laufzeiten der Maschinen:
A: 170 Stunden, B: 150 Stunden, C: 180 Stunden
 - Produkt 1: eine Stunde Maschine A, eine Stunde Maschine B
 - Produkt 2: zwei Stunden Maschine A, eine Stunde Maschine B und drei Stunden Maschine C
- Maximiere Profit, wenn 300 Euro mit Produkt 1 und 500 Euro mit Produkt 2 verdient wird,

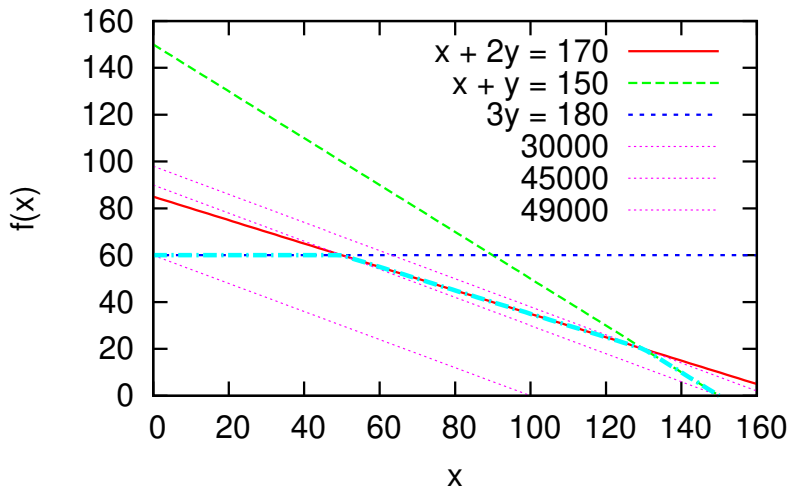
$$300 \cdot x_1 + 500 \cdot x_2 \rightarrow \max$$

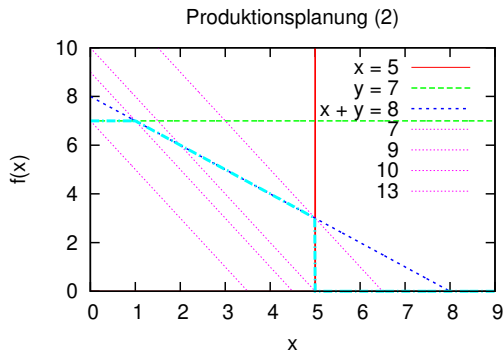
ohne dabei die Maschinenlaufzeiten zu überschreiten:

$$\begin{array}{rclcl} x_1 & + & 2 \cdot x_2 & \leq & 170 & \iff & x_2 & \leq & \frac{170-x_1}{2} \\ x_1 & + & x_2 & \leq & 150 & \iff & x_2 & \leq & 150 - x_1 \\ & & 3 \cdot x_2 & \leq & 180 & \iff & x_2 & \leq & \frac{180}{3} = 60 \end{array}$$

⁹http://de.wikipedia.org/wiki/Lineare_Optimierung

Produktionsplanung





Maximiere $2x_1 + x_2$

Randbedingungen:

$$0 \leq x_1 \leq 5$$

$$0 \leq x_2 \leq 7$$

$$0 \leq x_1 + x_2 \leq 8$$

Der optimale Wert liegt auf einem Eckpunkt des Polygons:

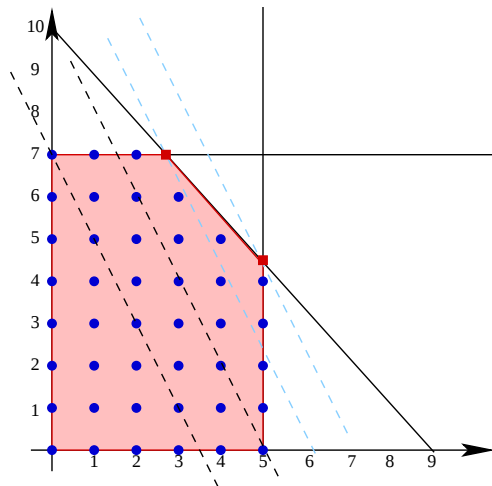
$$2x_1 + x_2 = 7 \iff x_2 = 7 - 2x_1$$

$$2x_1 + x_2 = 9 \iff x_2 = 9 - 2x_1$$

$$2x_1 + x_2 = 10 \iff x_2 = 10 - 2x_1$$

$$2x_1 + x_2 = 13 \iff x_2 = 13 - 2x_1$$

Verlangen wir, dass die Lösungen ganzzahlig sind, dann muss der optimale Wert nicht mehr auf einem Eckpunkt des Polygons liegen.



Maximiere $2x_1 + x_2$

Randbedingungen:

$$x_1 \leq 5 \text{ und } x_1 \in \mathbb{N}$$

$$x_2 \leq 7 \text{ und } x_2 \in \mathbb{N}$$

$$10x_1 + 9x_2 \leq 90$$

Viele kombinatorische Probleme können als Integer Programm formuliert werden.

Beispiel: 0/1-Rucksackproblem mit Profiten $p_1, \dots, p_n$ und den Gewichten $w_1, \dots, w_n$ der n Objekte sowie dem nicht zu überschreitenden Gesamtgewicht G :

$$x_j = \begin{cases} 1 & \text{das Objekt } j \text{ ist im Rucksack} \\ 0 & \text{sonst} \end{cases}$$

Binary Integer Program: $x_j \in \{0, 1\} \forall j$

$$\begin{array}{ll} \max & \sum_{j=1}^n p_j x_j \quad \rightarrow \text{Profit über alle Objekte im Rucksack} \\ \text{s.t.} & \sum_{j=1}^n w_j x_j \leq G \quad \rightarrow \text{Gesamtgewicht nicht überschreiten} \end{array}$$

Beispiel: Vertex Cover für Graph $G = (V, E)$ berechnen:

$$x_v = \begin{cases} 1 & \text{der Knoten } v \text{ ist im Vertex Cover} \\ 0 & \text{sonst} \end{cases}$$

Binary Integer Program: $x_v \in \{0, 1\} \quad \forall v \in V$

- Minimiere die Anzahl der Knoten im Vertex Cover:

$$\min \sum_{v \in V} x_v$$

- Randbedingung: $x_u + x_v \geq 1 \quad \forall \{u, v\} \in E$

Beispiel: Tourenplanung – Traveling Salesperson Problem für den Graph $G = (V, E, c)$ mit Kostenfunktion c :

$$x_{ij} = \begin{cases} 1 & \text{die Kante } (i, j) \text{ ist in der Tour} \\ 0 & \text{sonst} \end{cases}$$

Binary Integer Program: $x_{ij} \in \{0, 1\} \quad \forall i, j$

$$\begin{array}{lll} \min & \sum_{i,j} c_{ij} x_{ij} & \rightarrow \text{Kosten über alle benutzten Kanten} \\ \text{s.t.} & \sum_j x_{ij} = 1 \quad \forall i & \rightarrow \text{Knoten } i \text{ genau einmal verlassen} \\ & \sum_i x_{ij} = 1 \quad \forall j & \rightarrow \text{Knoten } j \text{ genau einmal besuchen} \end{array}$$

TSP mittels Integer Programming

Mit diesen Beschränkungen sind auch Teiltouren möglich, die keine Lösung des Problems beschreiben.

Wir fügen noch die Kurzzzyklus-Ungleichungen oder auch Subtour-Eliminationsbedingungen ein.

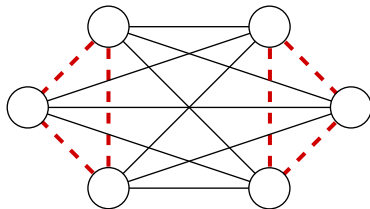
Sei S eine nicht-leere echte Teilmenge von $\{1, \dots, n\}$.

- Jede Tour muss S mindestens einmal verlassen:

$$\sum_{i \in S} \sum_{j \notin S} x_{ij} \geq 1$$

- Oder: In jeder Tour kann es nicht mehr als $|S| - 1$ Kanten geben, die Knoten in S verbinden, sonst gibt es eine Teiltour:

$$\sum_{i \in S} \sum_{j \in S} x_{ij} \leq |S| - 1$$



Lineare Programmierung:

gegeben: Matrix $A \in \mathbb{R}^{m,n}$ und zwei Vektoren $b \in \mathbb{R}^m$ und $c \in \mathbb{R}^n$

zulässige Lösung: Vektor $x \in \mathbb{R}^n$, der lineare Bedingungen erfüllt:

$$\begin{array}{ccccccc} a_{11}x_1 & + & a_{12}x_2 & + & \dots & + & a_{1n}x_n & \leq & b_1 \\ a_{21}x_1 & + & a_{22}x_2 & + & \dots & + & a_{2n}x_n & \leq & b_2 \\ \vdots & & \vdots & & & & \vdots & & \vdots \\ a_{m1}x_1 & + & a_{m2}x_1 & + & \dots & + & a_{mn}x_n & \leq & b_m \end{array}$$

Ziel: finde zulässigen Vektor x , der das Skalarprodukt

$$c^T x = c_1 x_1 + \dots + c_n x_n$$

maximiert. Dieses Optimierungsproblem wird oft abkürzend als

$$\max\{c^T x \mid Ax \leq b, x \geq 0\}$$

geschrieben, wobei die Bedingungen $Ax \leq b$ und $x \geq 0$ komponentenweise zu verstehen sind.

Überführen in Standardform:

- Minimieren statt Maximieren:

$$\min\{c^T x \mid Ax \leq b, x \geq 0\} \iff \max\{-c^T x \mid Ax \leq b, x \geq 0\}$$

- Größer-als statt Kleiner-als:

$$a_{i1}x_1 + \dots + a_{in}x_n \geq b_i \iff -a_{i1}x_1 - \dots - a_{in}x_n \leq -b_i$$

- Gleichheit statt Kleiner-gleich:

$$a_{i1}x_1 + \dots + a_{in}x_n = b_i$$

$$\iff a_{i1}x_1 + \dots + a_{in}x_n \geq b_i \wedge a_{i1}x_1 + \dots + a_{in}x_n \leq b_i$$

- Nichtnegativität der Variablen: ersetze x_i durch $x'_i - x''_i$ mit $x'_i, x''_i \geq 0$

F. Gurski: Efficient Binary Linear Programming Formulations for Boolean Functions. Statistics, Optimization and Information Computing, Vol 2, 274 – 279, 2014.

- Every conjunction $x_1 \wedge x_2$ can be realized by introducing a new variable x_3 and three conditions

$$x_3 - x_1 \leq 0, \quad x_3 - x_2 \leq 0, \quad x_1 + x_2 - x_3 \leq 1$$

such that finally $x_3 = x_1 \wedge x_2$.

- Every n -ary disjunction $x_1 \vee x_2 \vee \dots \vee x_n$ can be realized by introducing a new variable x_{n+1} and $n + 1$ conditions

$$x_1 - x_{n+1} \leq 0 \quad \dots \quad x_n - x_{n+1} \leq 0$$

and

$$x_1 + x_2 + \dots + x_n - x_{n+1} \geq 0$$

such that finally $x_{n+1} = x_1 \vee x_2 \vee \dots \vee x_n$.

Die innere Punkte-Methode¹⁰ oder die Ellipsoid-Methode liefern eine Lösung der Linearen Programmierung in polynomieller Zeit.

Probleme:

- Oft müssen die Variablen ganzzahlig sein.
 - Es können nicht 3.7 Flugzeuge gebaut werden.
 - Eine Maschine wird eingesetzt oder nicht.
- TSP: $O(2^{|E|})$ viele Subtouren müssen durch Eliminationsbedingungen ausgeschlossen werden. (Subtouren sinnvoll?)

Relaxation (Aufheben gewisser Einschränkungen):

- Lasse Ganzzahligkeitsbedingungen weg und berechne eine Lösung. Diese Lösung ist Schranke für eigentliches Problem.
- Nutze diese Schranke bei Branch-and-Bound-Verfahren, um den Suchraum zu begrenzen.

¹⁰N.K. Karmarkar: A new polynomial-time algorithm for linear programming. *Combinatorica*, 1984.

Mittels GLPK¹¹ (GNU Linear Programming Kit) können wir bspw. das n-Damen-Problem¹² lösen:

```
# size of the chess board
param n, integer, > 0, default 8;
# x[i,j] = 1 means a queen is placed in square [i,j]
var x{1..n, 1..n}, binary;

# at most one queen in each row, col, diagonal
subject to r{i in 1..n}: sum{j in 1..n} x[i,j] <= 1;
subject to c{j in 1..n}: sum{i in 1..n} x[i,j] <= 1;
s.t. diagl{k in 2-n..n-2}:
    sum{i in 1..n, j in 1..n: i-j == k} x[i,j] <= 1;
s.t. diagr{k in 3..n+n-1}:
    sum{i in 1..n, j in 1..n: i+j == k} x[i,j] <= 1;
```

¹¹<http://www.gnu.org/software/glpk/>

¹²</usr/share/doc/glpk-utils/examples/queens.mod>

```
# objective is to place as many queens as possible
maximize obj: sum{i in 1..n, j in 1..n} x[i,j];

# solve the problem
solve;

# print optimal solution
for {i in 1..n} {
    for {j in 1..n}
        printf " %s", if x[i,j] then "Q" else ".";
    printf("\n");
}

end;
# Written in GNU MathProg by Andrew Makhorin
```

Übung: Starten Sie obiges Programm mittels `glpsol --math queens.mod` für verschiedene Größen des Schachbretts.

GLPK enthält unter `/usr/share/doc/glpk-utils/examples/` Beispielprogramme für viele weitere Probleme:

- `sudoku.mod` ist ein Sudoku-Solver für das klassische Spiel der Größe 9×9 .
- `tsp.mod.gz` enthält Code zum Lösen des Travelling Salesperson Problems und einige kleinere Beispielgraphen.

Außerdem finden Sie unter `/usr/share/doc/glpk-doc/` das Reference Manual `glpk.pdf` zum GNU Linear Programming Kit sowie die Beschreibung `gmp1.pdf` der Modeling Language GNU MathProg.

Applegate, Bixby, Chvátal und Cook (1990): Beginn der Entwicklung des Programms Concorde, das an sämtlichen Lösungsrekorden der letzten Jahre beteiligt war.

Reinelt (1991): TSPLIB, Sammlung standardisierter Testinstanzen mit unterschiedlichem Schwierigkeitsgrad, anhand derer verschiedene Forschergruppen ihre Resultate vergleichen konnten.

Cook et al. (2006): beweisbar kürzeste Tour durch 85.900 Knoten eines Layoutproblems für integrierte Schaltkreise, was die bislang größte optimal gelöste TSPLIB-Instanz ist.¹³

Für Instanzen mit mehreren Millionen Städten konnten sie mit Hilfe zusätzlicher Dekompositionstechniken Touren bestimmen, deren Länge beweisbar weniger als 1% vom Optimum entfernt liegt.

¹³Applegate, Bixby, Chvátal, Cook: The Traveling Salesman Problem. A Computational Study. Princeton University Press, 2007.

¹⁴http://de.wikipedia.org/wiki/Problem_des_Handlungsreisenden

Ist es Zufall, dass der optimale Zielwert bei Linearer Programmierung auf einem Eckpunkt des Polygons liegt?

Beispiel:

$$\begin{array}{llllll} \text{maximiere} & 300 \cdot x_1 & + & 500 \cdot x_2 & & \\ \text{s.t.} & x_1 & + & 2 \cdot x_2 & \leq & 170 \quad (1) \\ & x_1 & + & x_2 & \leq & 150 \quad (2) \\ & & & 3 \cdot x_2 & \leq & 180 \quad (3) \end{array}$$

Wie sieht eine obere Schranke für den Zielwert aus?

$$\begin{array}{ll} 500 \cdot (2) & \rightarrow 500x_1 + 500x_2 \leq 75000 \\ 300 \cdot (1) & \rightarrow 300x_1 + 600x_2 \leq 51000 \\ 200 \cdot (1) + 100 \cdot (2) & \rightarrow 300x_1 + 500x_2 \leq 49000 \end{array}$$

Die *Zielfunktion* kann also maximal den Wert 49000 annehmen!

Wir haben eine Linearkombination der Ungleichungen $Ax \leq b$ gebildet, die gerade der ursprünglichen Zielfunktion entspricht, oder zumindest (komponentenweise) größer als die ursprüngliche Zielfunktion ist, also:

$$\lambda^T \cdot A \geq c^T$$

Da wir eine möglichst kleine obere Schranke für das ursprüngliche Problem haben möchten, müssen wir die neue Zielfunktion $\lambda^T \cdot b$ minimieren!

ursprüngliches Problem:

$$\begin{array}{ll} \text{maximiere} & c^T \cdot x \\ \text{s.t.} & A \cdot x \leq b \\ & x \geq 0 \end{array}$$

duales Problem:

$$\begin{array}{ll} \text{minimiere} & \lambda^T \cdot b \\ \text{s.t.} & \lambda^T \cdot A \geq c^T \\ & \lambda \geq 0 \end{array}$$

Die beiden Zielwerte sind gleich, wenn beide LP's eine Lösung haben.

Beispiel:

de.wikipedia.org/wiki/Simplex-Verfahren

$$\begin{array}{llllll} \text{maximiere} & z = & 300 \cdot x_1 & + & 500 \cdot x_2 & \\ \text{s.t.} & & x_1 & + & 2 \cdot x_2 & \leq 170 \\ & & x_1 & + & & x_2 \leq 150 \\ & & & & 3 \cdot x_2 & \leq 180 \\ & & x_1 \geq 0, & x_2 \geq 0 & & \end{array}$$

Zunächst überführen wir das gegebene lineare Programm durch Einführen von Schlupf-Variablen (engl. slack variables) in ein Gleichungssystem.

$$\begin{array}{llllllll} \text{maximiere} & z = & 300 \cdot x_1 & + & 500 \cdot x_2 & & & \\ \text{s.t.} & & x_1 & + & 2 \cdot x_2 & + & y_A & = 170 \\ & & x_1 & + & & x_2 & + & y_B = 150 \\ & & & & 3 \cdot x_2 & + & y_C & = 180 \\ & & x_1 \geq 0, & x_2 \geq 0, & y_A \geq 0, & y_B \geq 0, & y_C \geq 0 & \end{array}$$

Das Gleichungssystem ist unterbestimmt, da wir zusätzlich zu den Variablen noch so viele Schlupfvariablen wie Gleichungen haben:

$$\begin{array}{llllllll} \text{maximiere} & z = & 300 \cdot x_1 & + & 500 \cdot x_2 & & & \\ \text{s.t.} & & x_1 & + & 2 \cdot x_2 & + & y_A & = & 170 \\ & & x_1 & + & & x_2 & + & y_B & = & 150 \\ & & & & 3 \cdot x_2 & + & y_C & = & 180 \end{array}$$

Wir schreiben dieses Gleichungssystem so um, dass die Schlupf-Variablen auf der linken Seite stehen.

$$\begin{array}{llllll} z & = & 0 & + & 300 \cdot x_1 & + & 500 \cdot x_2 \\ y_A & = & 170 & - & x_1 & - & 2 \cdot x_2 \\ y_B & = & 150 & - & x_1 & - & x_2 \\ y_C & = & 180 & - & & - & 3 \cdot x_2 \end{array}$$

Eine initiale Lösung des Gleichungssystems kann man direkt ablesen: $x_1 = 0$, $x_2 = 0$, $y_A = 170$, $y_B = 150$ und $y_C = 180$.

Diese gefundene Lösung versuchen wir schrittweise zu verbessern, indem die Werte der Variablen x_1 und x_2 erhöht werden.

zunächst: Wie groß darf x_1 bzw. x_2 werden, ohne dass eine Restriktion verletzt wird?

- x_1 : da $x_2 = 0$ ist, erhalten wir folgende Gleichungen

$$0 \leq y_A = 170 - 1 \cdot x_1 \Rightarrow x_1 \leq 170$$

$$0 \leq y_B = 150 - 1 \cdot x_1 \Rightarrow x_1 \leq 150$$

- x_2 : da $x_1 = 0$ ist, erhalten wir folgende Gleichungen

$$0 \leq y_A = 170 - 2 \cdot x_2 \Rightarrow x_2 \leq 170/2$$

$$0 \leq y_B = 150 - 1 \cdot x_2 \Rightarrow x_2 \leq 150$$

$$0 \leq y_C = 180 - 3 \cdot x_2 \Rightarrow x_2 \leq 180/3$$

Wert der Zielfunktion $z = 300 \cdot x_1 + 500 \cdot x_2$ ist 45.000 für $x_1 = 150$ bzw. 30.000 für $x_2 = 60$. Also tauschen wir x_1 und y_B .

Da der Wert y_B auf 0 sinkt, wenn wir $x_1 = 150$ setzen, schreiben wir die entsprechende Zeile um, so dass x_1 links steht

$$y_B = 150 - x_1 - x_2 \iff x_1 = 150 - y_B - x_2$$

und ersetzen x_1 in obigen Gleichungen:

$$\begin{aligned} z &= 0 + 300 \cdot (150 - y_B - x_2) + 500 \cdot x_2 \\ y_A &= 170 - 1 \cdot (150 - y_B - x_2) - 2 \cdot x_2 \\ y_C &= 180 - 0 \cdot (150 - y_B - x_2) - 3 \cdot x_2 \end{aligned}$$

Als neues Gleichungssystem erhalten wir:

$$\begin{aligned} z &= 45.000 - 300 \cdot y_B + 200 \cdot x_2 \\ y_A &= 20 + 1 \cdot y_B - 1 \cdot x_2 \\ x_1 &= 150 - 1 \cdot y_B - 1 \cdot x_2 \\ y_C &= 180 + 0 \cdot y_B - 3 \cdot x_2 \end{aligned}$$

Links stehen die Variablen mit Wert ungleich 0.

Da die jetzige Zielfunktion noch einen positiven Koeffizienten enthält, kann die Lösung noch verbessert werden. Wir führen also einen weiteren Schritt durch, bei dem jetzt nur noch der Wert der Variablen x_2 geändert werden kann. Da $y_B = 0$ ist, gilt:

$$0 \leq y_A = 20 - 1 \cdot x_2 \Rightarrow x_2 \leq 20$$

$$0 \leq x_1 = 150 - 1 \cdot x_2 \Rightarrow x_2 \leq 150$$

$$0 \leq y_C = 180 - 3 \cdot x_2 \Rightarrow x_2 \leq 180/3$$

Wir setzen $x_2 = 20$ und y_A wird zu 0, daher schreiben wir die entsprechende Zeile um

$$y_A = 20 + 1 \cdot y_B - 1 \cdot x_2 \iff x_2 = 20 - y_A + y_B$$

und setzen dann x_2 in die Gleichungen ein:

$$z = 45.000 - 300 \cdot y_B + 200 \cdot (20 - y_A + y_B)$$

$$x_1 = 150 - 1 \cdot y_B - 1 \cdot (20 - y_A + y_B)$$

$$y_C = 180 + 0 \cdot y_B - 3 \cdot (20 - y_A + y_B)$$

Wir erhalten so schließlich:

$$\begin{array}{rclclcl} z & = & 49.000 & - & 100 \cdot y_B & - & 200 \cdot y_A \\ x_2 & = & 20 & + & 1 \cdot y_B & - & 1 \cdot y_A \\ x_1 & = & 130 & - & 2 \cdot y_B & + & 1 \cdot y_A \\ y_C & = & 120 & - & 3 \cdot y_B & + & 3 \cdot y_A \end{array}$$

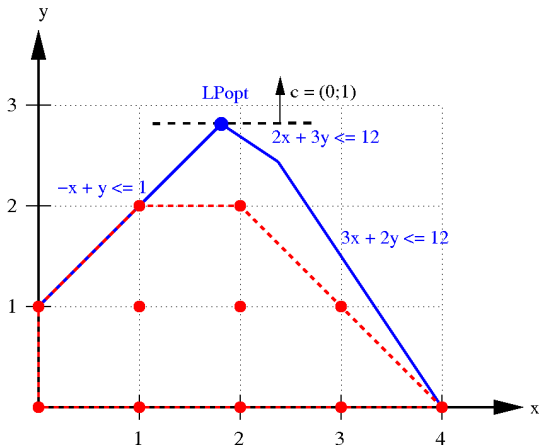
Auch hier stehen links die Variablen mit Wert ungleich 0. Da die Zielfunktion keine positiven Koeffizienten mehr enthält, ist eine optimale Lösung erreicht.

Als Ergebnis stellen wir fest:

- Die Maschinen *A* und *B* sind voll ausgelastet ($y_A = y_B = 0$),
- der Profit beträgt 49.000,
- Maschine *C* könnte noch 120 Stunden genutzt werden.

Integer Programming mittels LP-Relaxierung

Wir lösen zunächst das gegebene Problem mittels Simplex-Verfahren, lassen also die Ganzzahligkeitsbedingungen weg.



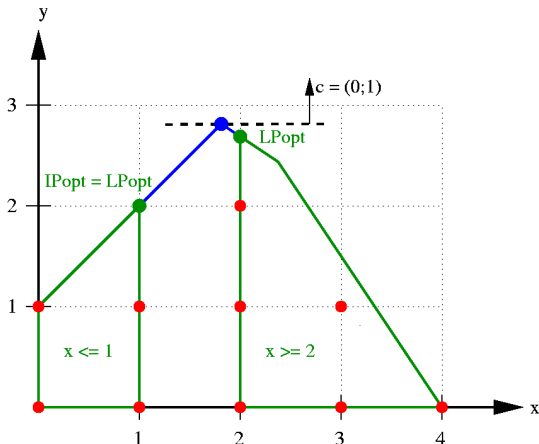
Quelle: de.wikipedia.org/wiki/Ganzzahlige_lineare_Optimierung

Die Lösung des relaxierten Problems liefert eine obere Schranke für den optimalen Wert der ganzzahligen Lösung, denn:

Jede zulässige Lösung des ganzzahligen Problems ist auch eine zulässige Lösung für das relaxierte Problem.

Der Wert der Lösung des relaxierten Problems wird *duale Schranke* genannt.

Für eine Variable x_j , die die Ganzzahligkeitsbedingung verletzt, fügen wir eine obere bzw. eine untere Schranke als Bedingung hinzu. Wir erhalten so zwei Teilprobleme, die wir wieder mit dem Simplex-Verfahren lösen.



Bei jeder Verzweigung (branch) wird also der Wertebereich einer Variablen eingeschränkt. Dies wird solange wiederholt, bis eine ganzzahlige Lösung für das Teilproblem gefunden wurde. Dann endet in diesem Zweig die Suche.

Der optimale Wert ergibt sich aus dem Maximum der gefundenen, ganzzahligen Lösungen der Teilprobleme.

Eine bereits gefundene ganzzahlige Lösung eines der bearbeiteten Zweige ist eine untere Schranke des Problems und wird *primale Schranke* genannt.

Ist die Lösung des relaxierten Teilproblems schlechter als die primale Schranke, dann wird die Bearbeitung des aktuellen Zweigs abgebrochen (bound).

Gap (absolut): Differenz zwischen dualer und primaler Schranke.

Gap (relativ): Gap (absolut) dividiert durch primale Schranke.

Sei x_i eine ganzzahlige Variable, die nach der LP-Relaxierung einen nicht-ganzzahligen Wert $\hat{x}_i$ hat. Dann fügen wir entweder neue Schranken hinzu oder wir ändern bereits hinzugefügte Schranken in der folgenden Weise:

- Falls x_i nicht begrenzt war, fügen wir $x_i \leq \lfloor \hat{x}_i \rfloor$ bzw. $\lceil \hat{x}_i \rceil \leq x_i$ als neue Schranke hinzu.
- Falls es bereits eine obere Schranke $x_i \leq c_2$ gibt, dann ändern wir diese Schranke in $x_i \leq \lfloor \hat{x}_i \rfloor$ bzw. $\lceil \hat{x}_i \rceil \leq x_i \leq c_2$.
- Falls es eine untere Schranke $c_1 \leq x_i$ gibt, dann ändern wir diese Schranke in $c_1 \leq x_i \leq \lfloor \hat{x}_i \rfloor$ bzw. $\lceil \hat{x}_i \rceil \leq x_i$.
- Falls die Variable bereits nach beiden Seiten $c_1 \leq x_i \leq c_2$ beschränkt ist, ändern wir die Schranke in $c_1 \leq x_i \leq \lfloor \hat{x}_i \rfloor$ bzw. $\lceil \hat{x}_i \rceil \leq x_i \leq c_2$.

In allen obigen Fällen erhalten wir zwei neue Teilprobleme.

Die Laufzeit des Branch-and-Bound-Verfahrens ist in der Regel zu groß, weil zu wenig vom Suchbaum abgeschnitten wird.

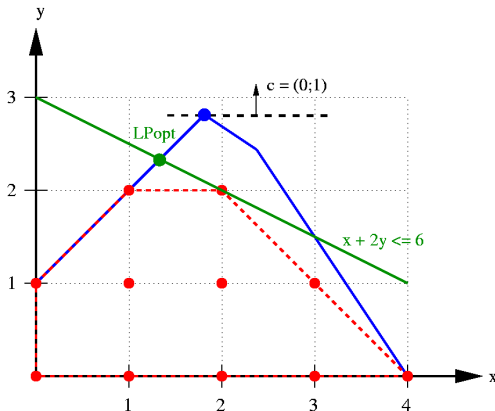
Gute IP-Löser verbessern daher die duale Schranke mittels Schnittebenenverfahren.

Schnittebene: zusätzliche Ungleichung, die von allen zulässigen Punkten des IPs erfüllt wird, aber nicht von der aktuellen LP-Lösung. (siehe nächste Folie)

Wird die Ungleichung dem LP hinzugefügt, muss beim erneuten Lösen eine andere Lösung herauskommen.

Integer Programming mittels LP-Relaxierung

Der Gap verkleinert sich, wenn eine kleinere, duale Schranke durch die Schnittebene (in der Grafik grün gezeichnet) gefunden wird.



Der Gap verkleinert sich auch dann, wenn eine bessere primale Schranke durch Lösen eines Teilproblems gefunden wird.

Lösen wir mittels Simplex-Verfahren das relaxierte lineare Programm aus dem obigen Bild.

$$\begin{array}{ll} \text{maximiere} & z = 1 \cdot x + 5 \cdot y \\ \text{s.t.} & -1 \cdot x + 1 \cdot y \leq 1 \\ & 3 \cdot x + 2 \cdot y \leq 12 \\ & 2 \cdot x + 3 \cdot y \leq 12 \\ & x \geq 0, y \geq 0 \end{array}$$

Nach Einfügen der Schlupfvariablen erhalten wir:

$$\begin{array}{ll} \text{maximiere} & z = 1 \cdot x + 5 \cdot y \\ \text{s.t.} & -1 \cdot x + 1 \cdot y + a = 1 \\ & 3 \cdot x + 2 \cdot y + b = 12 \\ & 2 \cdot x + 3 \cdot y + c = 12 \\ & x \geq 0, y \geq 0, a \geq 0, b \geq 0, c \geq 0 \end{array}$$

Als initiale Simplex-Tabelle ergibt sich:

$$z = 0 + 1 \cdot x + 5 \cdot y \quad (0)$$

$$a = 1 + 1 \cdot x - 1 \cdot y \quad (1)$$

$$b = 12 - 3 \cdot x - 2 \cdot y \quad (2)$$

$$c = 12 - 2 \cdot x - 3 \cdot y \quad (3)$$

Eine initiale Lösung ergibt sich für $x = 0$ und $y = 0$. Anhand der Nebenbedingungen ermitteln wir, wie groß y werden darf:

$$0 \leq a = 1 - 1 \cdot y \Rightarrow y \leq 1$$

$$0 \leq b = 12 - 2 \cdot y \Rightarrow y \leq 6$$

$$0 \leq c = 12 - 3 \cdot y \Rightarrow y \leq 4$$

Wir setzen $y = 1$, formen Gleichung (1) nach y um, setzen $y = 1 + x - a$ in (0), (2) und (3) ein und erhalten:

$$\begin{aligned} z &= 0 + 1 \cdot x + 5 \cdot (1 + x - a) \\ &= 5 + 6x - 5a \end{aligned} \quad (0')$$

$$\begin{aligned} y &= 1 + 1 \cdot x - 1 \cdot a \\ &= 1 + x - a \end{aligned} \quad (1')$$

$$\begin{aligned} b &= 12 - 3 \cdot x - 2 \cdot (1 + x - a) \\ &= 10 - 5x + 2a \end{aligned} \quad (2')$$

$$\begin{aligned} c &= 12 - 2 \cdot x - 3 \cdot (1 + x - a) \\ &= 9 - 5x + 3a \end{aligned} \quad (3')$$

Wir ermitteln nun, wie groß x werden darf:

$$0 \leq y = 1 + 1 \cdot x \Rightarrow x \leq \infty$$

$$0 \leq b = 10 - 5 \cdot x \Rightarrow x \leq 2$$

$$0 \leq c = 9 - 5 \cdot x \Rightarrow x \leq \frac{9}{5}$$

Wir setzen $x = 9/5$, formen Gleichung (3') nach x um, setzen $x = 9/5 + 3/5a - 1/5c$ in (0'), (1') und (2') ein und erhalten:

$$\begin{aligned}z &= 5 + 6 \cdot (9/5 + 3/5a - 1/5c) - 5 \cdot a \\&= 15 \frac{4}{5} - 1 \frac{2}{5}a - 1 \frac{1}{5}c \\y &= 1 + 1 \cdot (9/5 + 3/5a - 1/5c) - 1 \cdot a \\&= 2 \frac{4}{5} - \frac{2}{5}a - \frac{1}{5}c \\b &= 10 - 5 \cdot (9/5 + 3/5a - 1/5c) + 2 \cdot a \\&= 1 - 3a + c \\x &= \frac{9}{5} + \frac{3}{5} \cdot a - \frac{1}{5} \cdot c \\&= 1 \frac{4}{5} + \frac{3}{5}a - \frac{1}{5}c\end{aligned}$$

Um Gomory-Cuts zu berechnen, trennen wir den ganzzahligen Teil vom nicht-ganzzahligen Teil einer der obigen Gleichungen:

$$\begin{aligned}y = 2 \frac{4}{5} - \frac{2}{5}a - \frac{1}{5}c &\iff y + \frac{2}{5}a + \frac{1}{5}c = 2 \frac{4}{5} \\&\iff (y - 2) + \frac{2}{5}a + \frac{1}{5}c = \frac{4}{5} \\&\Rightarrow \frac{2}{5}a + \frac{1}{5}c \geq \frac{4}{5}\end{aligned}$$

Die letzte Zeile ist ein Gomory-Cut.

$$2/5a + 1/5c \geq 4/5 \iff 2a + 1c \geq 4$$

Wollen wir diesen Cut in unser Bild einzeichnen, müssen wir die zusätzliche Nebenbedingung bezüglich der x, y -Werte umrechnen:

$$\begin{aligned} -x + 1y + a &= 1 &\iff -2x + 2y + 2a &= 2 \\ 2x + 3y + c &= 12 \end{aligned}$$

Wenn wir diese Gleichungen addieren und den Gomory-Cut einsetzen, erhalten wir:

$$5y + 2a + c = 14 \iff 5y \leq 10 \iff y \leq 2$$

Durch diesen Cut schneiden wir die aktuelle Lösung des LP-Problems ab, behalten aber alle ganzzahligen Lösungen.

Übung: Lösen Sie das folgende lineare Programm mittels des Simplex-Verfahrens und bestimmen Sie einen Gomory-Cut.

$$\begin{array}{ll} \text{maximiere} & y \\ \text{s.t.} & -x + y \leq 1 \\ & x + 2y \leq 10 \end{array}$$

Zeichnen Sie den Gomory-Cut sowie die obigen Geraden in ein Koordinatensystem ein.

```
#include <glpk.h>
.....
glp_prob *lp = glp_create_prob();
glp_read_mps(lp, GLP_MPS_FILE, 0, "filename");
glp_term_out(GLP_ON);

glp_adv_basis(lp, 0);
int rc = glp_simplex(lp, 0);
if (rc != 0 || glp_get_status(lp) != GLP_OPT) {
    cout << "error" << endl; return 1;
}
cout << "lp: " << glp_get_obj_val(lp) << endl;

rc = glp_intopt(lp, 0);
if (rc == 0)
    cout<< "opt: " << glp_mip_obj_val(lp)<< endl;
else cout "error" << endl;
glp_delete_prob(lp);
```

MIPLIB 2003:

<http://miplib.zib.de/miplib2003/>

- electronically available library of both pure and mixed integer programs
- has become a standard test set used to compare the performance of mixed integer optimizers

Übung: Starten Sie `glpsol` für verschiedene Instanzen der MIPLIB. Testen Sie verschiedene Heuristiken für Node-Selection und Branching.

Auszug aus `glpsol --help`:

```
--first    branch: first integer variable
--last     branch: last integer variable
--mostf    branch: most fractional variable
--drtom     branch: heuristic by Driebeck and Tomlin(*)
--pcost     branch: hybrid pseudocost heuristic
```

Auszug aus glpsol --help: (Fortsetzung)

--dfs	backtrack: depth first search
--bfs	backtrack: breadth first search
--bestp	backtrack: best projection heuristic
--bestb	backtrack: best local bound(*)
--gomory	generate Gomory's mixed integer cuts
--mir	generate Mixed Integer Rounding cuts
--cover	generate mixed cover cuts
--clique	generate clique cuts
--cuts	generate all cuts above

Branching Rules: Let x_i be a variable with fractional LP value.

- *Most Fractional Branching* choose the variable with fractional part closest to 0.5

$$\text{score}(x_i) := 0.5 - | \tilde{x}_i - \lfloor \tilde{x}_i \rfloor - 0.5 |$$

Wähle x_j mit $\text{score}(x_j)$ maximal. $\text{score}(7.9) = 0.1$,
 $\text{score}(3.6) = 0.4$

- *Least Fractional Branching* $\text{score}(x_i) := | \tilde{x}_i - \lfloor \tilde{x}_i \rfloor - 0.5 |$

Wähle x_j mit $\text{score}(x_j)$ maximal. $\text{score}(7.9) = 0.4$,
 $\text{score}(3.6) = 0.1$

- *Full Strong Branching* test which of the fractional candidates gives the best progress

Δ_l change in objective value of left branch on x_i

Δ_r change in objective value of right branch on x_i

$$\text{score}(x_i) := (1 - \mu) \cdot \min\{\Delta_l, \Delta_r\} + \mu \cdot \max\{\Delta_l, \Delta_r\}$$

- *Pseudo-Cost Branching*

```
#include <iostream>
#include <list>
#include <cmath>
#include <cfloat>
#include <glpk.h>
#include <unistd.h>

using namespace std;

const double _eps = 1e-8;

list<int> getFractionalVariables(glp_prob *p);
int getBranchVariable(glp_prob *p);
double getScore(glp_prob *p, int var);
pair<glp_prob *, glp_prob *>
    getBranches(glp_prob *job, int i);
```

```
int getBranchVariable(glp_prob *p) {  
    double val = 0.0;  
    int var = -1;  
    list<int>::iterator it;  
    list<int> vars = getFractionalVariables(p);  
  
    for (it = vars.begin(); it != vars.end();  
         it++) {  
        double score = getScore(p, *it);  
  
        if (score > val) {  
            val = score;  
            var = *it;  
        }  
    }  
    return var;  
}
```

```
list<int> getFractionalVariables(glp_prob *p) {  
    list<int> vars;  
    int cols = glp_get_num_cols(p);  
  
    for (int i = 1; i <= cols; i++) {  
        int knd = glp_get_col_kind(p, i);  
  
        if (glp_get_col_stat(p, i) == GLP_BS &&  
            (knd == GLP_IV || knd == GLP_BV)) {  
            double xi = glp_get_col_prim(p, i);  
            double rVal = round(xi);  
            double diff = fabs(xi - rVal);  
  
            if (diff >= _eps) vars.push_back(i);  
        }  
    }  
    return vars;  
}
```

```
double getScore(glp_prob *p, int var) {  
    // -----  
    // Most Fractional Branch  
    // -----  
    double val = glp_get_col_prim(p, var);  
  
    return 0.5 - fabs(val - floor(val) - 0.5);  
}
```

```
pair<glp_prob *, glp_prob *>
    getBranches(glp_prob *job, int i) {
    double val = glp_get_col_prim(job, i);
    double ub = floor(val);
    double lb = ub + 1.0;

    glp_prob *lp = glp_create_prob();
    glp_copy_prob(lp, job, GLP_OFF);

    glp_prob *rp = glp_create_prob();
    glp_copy_prob(rp, job, GLP_OFF);

    int type = glp_get_col_type(job, i);
    double lower = glp_get_col_lb(job, i);
    double upper = glp_get_col_ub(job, i);
```

```
int lType , rType ;

if (type == GLP_DB) {
    lType = GLP_DB ;
    rType = GLP_DB ;
} else if (type == GLP_LO) {
    lType = GLP_DB ;
    rType = GLP_LO ;
} else if (type == GLP_UP) {
    lType = GLP_UP ;
    rType = GLP_DB ;
} else {
    lType = GLP_UP ;
    rType = GLP_LO ;
}
```

```
    if (lower >= ub) {  
        ub = lower;  
        lType = GLP_FX;  
    }  
    glp_set_col_bnds(lp, i, lType, lower, ub);  
  
    if (lb >= upper) {  
        lb = upper;  
        rType = GLP_FX;  
    }  
    glp_set_col_bnds(rp, i, rType, lb, upper);  
  
    return pair<glp_prob *, glp_prob *>(lp, rp);  
}
```

```
int main(int argc, char *argv[]) {
    if (argc != 2) {
        cout << "usage: " << argv[0]
              << " filename" << endl;
        return 1;
    }

    double opt;
    char *file = argv[1];

    glp_prob *lp = glp_create_prob();
    if (glp_read_mps(lp, GLP_MPS_FILE, 0, file)){
        cout << file << " not found" << endl;
        return 1;
    }

    // disable terminal output
    glp_term_out(GLP_OFF);
}
```

```
int objDir = glp_get_obj_dir(lp);

if (objDir == GLP_MAX)
    opt = -DBL_MAX;
else opt = DBL_MAX;

glp_adv_basis(lp, 0);
int rc = glp_simplex(lp, 0);    // lp relaxation

if (rc != 0 || glp_get_status(lp) != GLP_OPT){
    cout << "error" << endl;
    return 1;
}
cout << "==> lp-relax: ";
if (objDir == GLP_MIN)
    cout << "lower bound ";
else cout << "upper bound ";
cout << glp_get_obj_val(lp) << " <==" << endl;
```

```
list<glp_prob *> pool;

pool.push_back(lp);
while ( ! pool.empty() ) {
    glp_prob *job = pool.front();
    pool.pop_front();

    glp_adv_basis(job, 0);
    rc = glp_simplex(job, 0);    // lp relaxation

    if (0 != rc ||
        glp_get_status(job) != GLP_OPT) {
        // no valid solution found
        // ==> continue with next job
        glp_delete_prob(job);
        continue;
    }
}
```

```
double val = glp_get_obj_val(job);
if ((objDir == GLP_MAX && val <= opt)
    || (objDir == GLP_MIN && val >= opt)) {
    // current solution is bad ==> cut
    glp_delete_prob(job);
    continue;
}

int i = getBranchVariable(job);

if (i == -1) {
    if ((objDir == GLP_MAX && val > opt)
        || (objDir == GLP_MIN && val < opt)){
        opt = val;
        cout << "solution: " << val << endl;
    }
    glp_delete_prob(job); continue;
}
```

```
pair<glp_prob *, glp_prob *> probs =  
    getBranches(job, i);  
  
    // Depth First Search  
    pool.push_front(probs.first);  
    pool.push_front(probs.second);  
  
    glp_delete_prob(job);  
}  
cout << ">>> value: " << opt << " <<<" << endl;  
  
return 0;  
}
```

Darstellen boolescher Logik mittels Integer Programmen

$y = \neg x$ wird zu $y = 1 - x$

x	$\neg x$	y	$y = \neg x$
0	1	0	0
0	1	1	1
1	0	0	1
1	0	1	0

$\leadsto$

x	$1 - x$	y	$y = 1 - x$
0	1	0	0
0	1	1	1
1	0	0	1
1	0	1	0

$x \Rightarrow y$ wird zu $x \leq y$

x	y	$x \Rightarrow y$
0	0	1
0	1	1
1	0	0
1	1	1

$\leadsto$

x	y	$x \leq y$
0	0	1
0	1	1
1	0	0
1	1	1

Darstellen boolescher Logik mittels Integer Programmen

$z = x \vee y$ wird zu $0 \leq 2 \cdot z - x - y \leq 1$

x	y	z	$z = x \vee y$
0	0	0	1
0	0	1	0
0	1	0	0
0	1	1	1
1	0	0	0
1	0	1	1
1	1	0	0
1	1	1	1

$\leadsto$

x	y	z	$0 \leq 2 \cdot z - x - y \leq 1$
0	0	0	1
0	0	1	0
0	1	0	0
0	1	1	1
1	0	0	0
1	0	1	1
1	1	0	0
1	1	1	1

allgemein: $y = x_1 \vee x_2 \vee \dots \vee x_n$ wird zu $0 \leq n \cdot y - \sum_i x_i \leq n - 1$

Darstellen boolescher Logik mittels Integer Programmen

$z = x \wedge y$ wird zu $0 \leq x + y - 2 \cdot z \leq 1$

x	y	z	$z = x \wedge y$
0	0	0	1
0	0	1	0
0	1	0	1
0	1	1	0
1	0	0	1
1	0	1	0
1	1	0	0
1	1	1	1

$\leadsto$

x	y	z	$0 \leq x + y - 2 \cdot z \leq 1$
0	0	0	1
0	0	1	0
0	1	0	1
0	1	1	0
1	0	0	1
1	0	1	0
1	1	0	0
1	1	1	1

allgemein: $y = x_1 \wedge x_2 \wedge \dots \wedge x_n$ wird zu $0 \leq \sum_i x_i - n \cdot y \leq n - 1$

Darstellen boolescher Logik mittels Integer Programmen

$$\begin{aligned}(x \wedge y) \Rightarrow z &\equiv \neg(x \wedge y) \vee z \\ &\equiv \neg x \vee \neg y \vee z\end{aligned}$$

wird zu $(1 - x) + (1 - y) + z \geq 1$

x	y	z	$(x \wedge y) \Rightarrow z$
0	0	0	1
0	0	1	1
0	1	0	1
0	1	1	1
1	0	0	1
1	0	1	1
1	1	0	0
1	1	1	1

$\leadsto$

x	y	z	$(1 - x) + (1 - y) + z \geq 1$
0	0	0	1
0	0	1	1
0	1	0	1
0	1	1	1
1	0	0	1
1	0	1	1
1	1	0	0
1	1	1	1

Darstellen boolescher Logik mittels Integer Programmen

$$\underbrace{\neg A \wedge \neg B \Rightarrow C = D}_{=:F} \sim \underbrace{C - D \leq A + B}_{=:X} \text{ und } \underbrace{-C + D \leq A + B}_{=:Y}$$

A	B	C	D	F	X	Y
0	0	0	0	1	1	1
0	0	0	1	0	1	0
0	0	1	0	0	0	1
0	0	1	1	1	1	1
0	1	0	0	1	1	1
0	1	0	1	1	1	1
0	1	1	0	1	1	1
0	1	1	1	1	1	1
1	0	0	0	1	1	1
1	0	0	1	1	1	1
⋮				⋮		

Use binary decision variables:

- $X_{i,t} = 1$ iff position i is empty after move t
- $Y_{i,j,t} = 1$ iff peg is moved from position i to j at time t

The objective is to maximize the number of empty holes in the final state:

$$\sum_i X_{i,T} \longrightarrow \max$$

T denotes the number of moves which depends on the number of pegs in the initial state and on the number of pegs in the final state, since after each move exactly one peg has been removed from the board.

We have to consider the following constraints:

- One jump per turn:

$$\sum_i \sum_{j \in D_i} Y_{i,j,t} = 1 \quad \forall t$$

- Jump to empty hole: $\bigvee_{j \in D_i} Y_{j,i,t} \Rightarrow X_{i,t}$

$$\sum_{j \in D_i} Y_{j,i,t} \leq X_{i,t} \quad \forall i, t$$

- Jump from occupied hole: $\bigvee_{j \in D_i} Y_{i,j,t} \Rightarrow \neg X_{i,t}$

$$\sum_{j \in D_i} Y_{i,j,t} \leq 1 - X_{i,t} \quad \forall i, t$$

D_i denotes the set of destination positions for a peg in position i .

further constraints:

- Jump over peg: $\bigvee_{(j,k) \in E_i} Y_{j,k,t} \Rightarrow \neg X_{i,t}$

$$\sum_{(j,k) \in E_i} Y_{j,k,t} \leq 1 - X_{i,t} \quad \forall i, t$$

- constraints to describe the initial and final state

E_i denotes the position pairs (j, k) such that a jump from position j over position i to position k is possible.

Further we have to ensure the balance equations:

$$X_{i,t} + \sum_{(j,k) \in E_i} Y_{j,k,t} + \sum_{j \in D_i} Y_{i,j,t} - \sum_{j \in D_i} Y_{j,i,t} = X_{i,t+1} \quad \forall i, t$$

The balance equation ensures:

- $X_{i,t} = X_{i,t+1}$ if position i is not involved in a move
- $X_{i,t} = 1$ and $X_{i,t+1} = 0$ if a move is made to i
- $X_{i,t} = 0$ and $X_{i,t+1} = 1$ if a move is made from i or over i

This IP model comes from:

G.W. DePuy and G.D. Taylor: Using board puzzles to teach operations research. *INFORMS Transactions on Education* 7(2), 160 – 171 (2007)

Planen mittels Integer Programming: Knight Tour

1				
	2			

1				
	2			
3				

...

1	24	13	18	7
14	19	8	23	12
9	2	xx	6	17
20	15	4	11	22
3	10	21	16	5

Use binary decision variables:

- $X_{i,t} = 1$ iff square i contains the knight at time t
- $Y_{ij,t} = 1$ iff at time t the knight is moved from square i to j

The objective is to maximize the number of moves:

$$\sum_i \sum_{j \in D_i} \sum_t Y_{ij,t} \longrightarrow \max$$

D_i denotes the set of feasible positions that can be reached by the knight from square i .

Planen mittels Integer Programming: Knight Tour

We have to consider the following constraints:

- the knight stands on one square per turn:

$$\sum_i X_{i,t} = 1 \quad \forall t$$

- the knight visits each square once over time:

$$\sum_t X_{i,t} = 1 \quad \forall i$$

- exactly one move is made per turn:

$$\sum_i \sum_{j \in D_i} Y_{i,j,t} = 1 \quad \forall t$$

- a move from square i requires the knight to be there:

$$\bigvee_{j \in D_i} Y_{i,j,t} \Rightarrow X_{i,t} \quad \rightsquigarrow \quad \sum_{j \in D_i} Y_{i,j,t} \leq X_{i,t} \quad \forall i, t$$

Further we have to ensure the balance equation:

$$X_{i,t} - \sum_{j \in D_i} Y_{i,j,t} + \sum_{j \in D_i} Y_{j,i,t} = X_{i,t+1} \quad \forall i, t$$

The balance equation ensures:

- $X_{i,t} = X_{i,t+1}$ if square i is not involved in a move
- $X_{i,t} = 1$ and $X_{i,t+1} = 0$ if a move is made from i
- $X_{i,t} = 0$ and $X_{i,t+1} = 1$ if a move is made to i

Knight-Tour problem is efficiently solvable, see

I. Parberry: An efficient algorithm for the knight's tour problem.
Discrete Applied Mathematics 73(3), 251 – 260 (1997)

Planen mittels Integer Programming: 15-Puzzle

2	6	3	2	6	3	2	6	3	2		3		2	3	
1	7	8	1	7	8	1		8	1	6	8	1	6	8	...
4	5		4		5	4	7	5	4	7	5	4	7	5	

Use binary decision variables:

- $X_{i,k,t} = 1$ iff square i contains number k at time t
number 0 denotes the missing tile
- $Y_{i,j,t} = 1$ iff at time t the numbers of the squares i and j exchange places

The objective is to minimize the number of moves:

$$\sum_i \sum_{j \in D_i} \sum_t Y_{i,j,t} \rightarrow \min$$

D_i denotes the set of destination positions if the empty tile is in square i . Value T denotes the number of moves which must be chosen large enough.

We have to consider the following constraints:

- exchange at most one pair of squares per turn:

$$\sum_i \sum_{j \in D_i} Y_{i,j,t} \leq 1 \quad \forall t$$

- each square contains exactly one number at each time:

$$\sum_k X_{i,k,t} = 1 \quad \forall i, t$$

- the puzzle contains each number exactly once per time:

$$\sum_i X_{i,k,t} = 1 \quad \forall k, t$$

Planen mittels Integer Programming: 15-Puzzle

- if empty tile is moved from square i to j
 - then square i holds 0 at time t :

$$Y_{i,j,t} \Rightarrow X_{i,0,t} \leadsto Y_{i,j,t} \leq X_{i,0,t} \quad \forall i, j \in D_i, t$$

- then square j holds 0 at time $t + 1$:

$$Y_{i,j,t} \Rightarrow X_{j,0,t+1} \leadsto Y_{i,j,t} \leq X_{j,0,t+1} \quad \forall i, j \in D_i, t$$

- and square j holds number k , then square i holds k afterwards:

$$Y_{i,j,t} \wedge X_{j,k,t} \Rightarrow X_{i,k,t+1}$$

$$1 - Y_{i,j,t} + 1 - X_{j,k,t} + X_{i,k,t+1} \geq 1 \quad \forall i, j \in D_i, t, k > 0$$

- if no exchange between squares i and j has been made, then $X_{i,k,t} = X_{i,k,t+1}$.

$$X_{i,k,t} - X_{i,k,t+1} \leq \sum_{j \in D_i} Y_{i,j,t} + \sum_{j \in D_i} Y_{j,i,t} \quad \forall i, k, t$$

$$-X_{i,k,t} + X_{i,k,t+1} \leq \sum_{j \in D_i} Y_{i,j,t} + \sum_{j \in D_i} Y_{j,i,t} \quad \forall i, k, t$$

Planen mittels Integer Programming: Rush-Hour



Use binary decision variables, where the marker denotes the upper left position of a vehicle:

- $X_{v,j,t} = 1$ iff marker of vehicle v is at position j at time t
- $Z_{v,j,t} = 1$ iff vehicle v occupies position j at time t .
- $Y_{v,i,j,t} = 1$ iff vehicle v moves marker from position i to position j at time t .

The objective is to minimize the number of moves:

$$\sum_v \sum_{i \in B_v} \sum_{j \in B_v} \sum_t Y_{v,i,j,t} \longrightarrow \min$$

B_v denotes the set of possible positions of marker of vehicle v .

The number of moves T must be chosen large enough to be able to solve the problem.

We have to consider the following constraints:

- at most one vehicle in a position:

$$\sum_v Z_{v,i,t} \leq 1 \quad \forall i, t$$

- move at most one vehicle per turn:

$$\sum_v \sum_{i \in B_v} \sum_{j \in B_v} Y_{v,i,j,t} \leq 1 \quad \forall t$$

- define positions occupied by vehicles:

$$L_v \cdot X_{v,i,t} \leq \sum_{m \in M_{v,i}} Z_{v,m,t} \quad \forall v, i, t$$

L_v denotes the length of vehicle v , and $M_{v,j}$ denotes the set of positions occupied by vehicle v if its marker is in position j .

Further constraints:

- take note of occupied positions:

$$Y_{v,i,j,t} \leq 1 - \sum_{w \neq v} Z_{w,p,t-1} \quad \forall v, i, j, t, p \in P_{v,i,j}$$

$P_{v,i,j}$ denotes the set of positions between i and j that vehicle v may reach.

- update marker for vehicles:

$$X_{v,i,t-1} - \sum_{j \in B_v} Y_{v,i,j,t} + \sum_{j \in B_v} Y_{v,j,i,t} = X_{v,i,t} \quad \forall v, i, t$$

Vorgehen bei gegebener Modellierung mittels STRIPS:

- Wird eine Aktion a durchgeführt, dann müssen vorher die PRE-Conditions und hinterher die ADD-Effects gelten, die DEL-Effects dürfen nicht mehr gelten:

$$\begin{aligned}a_t &\rightarrow p_t && \forall p \in pre(a) \\a_t &\rightarrow p_{t+1} && \forall p \in add(a) \\a_t &\rightarrow \neg p_{t+1} && \forall p \in del(a)\end{aligned}$$

- Parallele, widersprüchliche Aktionen werden ausgeschlossen:

$$a_t \rightarrow \neg b_t \quad \forall a, b : del(a) \cap (pre(b) \cup add(b)) \neq \emptyset$$

- Bei einem Wechsel des Zustands muss eine entsprechende Aktion durchgeführt worden sein:

$$\begin{aligned}(\neg p_t \wedge p_{t+1}) &\rightarrow \bigvee_{p \in add(a)} a_t \\(p_t \wedge \neg p_{t+1}) &\rightarrow \bigvee_{p \in del(a)} a_t\end{aligned}$$

`unstack(X,Y):`

- PRE: `upon(X,Y)`, `clear(X)`, `empty()`
- ADD: `clear(Y)`, `hold(X)`
- DEL: `upon(X,Y)`, `clear(X)`, `empty()`

Constraints im linearen Programm:

$$\begin{aligned}unstack[x, y, t] &\leq upon[x, y, t] \\unstack[x, y, t] &\leq clear[x, t] \\unstack[x, y, t] &\leq empty[t] \\unstack[x, y, t] &\leq clear[y, t + 1] \\unstack[x, y, t] &\leq hold[x, t + 1] \\unstack[x, y, t] &\leq 1 - upon[x, y, t + 1] \\unstack[x, y, t] &\leq 1 - clear[x, t + 1] \\unstack[x, y, t] &\leq 1 - empty[t + 1]\end{aligned}$$

`stack(X,Y):`

- PRE: `clear(Y)`, `hold(X)`
- ADD: `upon(X,Y)`, `clear(X)`, `empty()`
- DEL: `clear(Y)`, `hold(X)`

Constraints im linearen Programm:

$$\begin{aligned}stack[x, y, t] &\leq clear[y, t] \\stack[x, y, t] &\leq hold[x, t] \\stack[x, y, t] &\leq upon[x, y, t + 1] \\stack[x, y, t] &\leq clear[x, t + 1] \\stack[x, y, t] &\leq empty[t + 1] \\stack[x, y, t] &\leq 1 - clear[y, t + 1] \\stack[x, y, t] &\leq 1 - hold[x, t + 1]\end{aligned}$$

`pickup(X)`:

- PRE: `onTable(X)`, `clear(X)`, `empty()`
- ADD: `hold(X)`
- DEL: `onTable(X)`, `clear(X)`, `empty()`

Constraints im linearen Programm:

$$\begin{aligned}pickup[x, t] &\leq onTable[x, t] \\pickup[x, t] &\leq clear[x, t] \\pickup[x, t] &\leq empty[t] \\pickup[x, t] &\leq hold[x, t + 1] \\pickup[x, t] &\leq 1 - onTable[x, t + 1] \\pickup[x, t] &\leq 1 - clear[x, t + 1] \\pickup[x, t] &\leq 1 - empty[t + 1]\end{aligned}$$

`putdown(X)`:

- PRE: `hold(X)`
- ADD: `onTable(X)`, `clear(X)`, `empty()`
- DEL: `hold(X)`

Constraints im linearen Programm:

$$\begin{aligned} \textit{putdown}[x, t] &\leq \textit{hold}[x, t] \\ \textit{putdown}[x, t] &\leq \textit{onTable}[x, t + 1] \\ \textit{putdown}[x, t] &\leq \textit{clear}[x, t + 1] \\ \textit{putdown}[x, t] &\leq \textit{empty}[t + 1] \\ \textit{putdown}[x, t] &\leq 1 - \textit{hold}[x, t + 1] \end{aligned}$$

Falls sich $upon(x, y)$ geändert hat:

$$\begin{aligned}1 - upon[x, y, t] + upon[x, y, t + 1] &\leq 1 + stack[x, y, t] \\ upon[x, y, t] + 1 - upon[x, y, t + 1] &\leq 1 + unstack[x, y, t]\end{aligned}$$

Falls sich $clear(x)$ geändert hat:

$$\begin{aligned}1 - clear[x, t] + clear[x, t + 1] \\ &\leq 1 + \sum_y unstack[y, x, t] + \sum_y stack[x, y, t] + putdown[x, t] \\ clear[x, t] + 1 - clear[x, t + 1] \\ &\leq 1 + \sum_y stack[y, x, t] + \sum_y unstack[x, y, t] + pickup[x, t]\end{aligned}$$

Falls sich $onTable(x)$ geändert hat:

$$\begin{aligned}1 - onTable[x, t] + onTable[x, t + 1] &\leq 1 + putdown[x, t] \\ onTable[x, t] + 1 - onTable[x, t + 1] &\leq 1 + pickup[x, t]\end{aligned}$$

Falls sich `empty()` geändert hat:

$$\begin{aligned} & 1 - \text{empty}[t] + \text{empty}[t + 1] \\ & \leq 1 + \sum_x \text{putdown}[x, t] + \sum_x \sum_y \text{stack}[x, y, t] \\ & \text{empty}[t] + 1 - \text{empty}[t + 1] \\ & \leq 1 + \sum_x \text{pickup}[x, t] + \sum_x \sum_y \text{unstack}[x, y, t] \end{aligned}$$

Falls sich `hold(x)` geändert hat:

$$\begin{aligned} & 1 - \text{hold}[x, t] + \text{hold}[x, t + 1] \\ & \leq 1 + \text{pickup}[x, t] + \sum_y \text{unstack}[x, y, t] \\ & \text{hold}[x, t] + 1 - \text{hold}[x, t + 1] \\ & \leq 1 + \text{putdown}[x, t] + \sum_y \text{stack}[x, y, t] \end{aligned}$$

Es darf nur eine Operation zur Zeit t ausgeführt werden:

$$\sum_x \sum_y (stack[x, y, t] + unstack[x, y, t]) \\ + \sum_x (pickup[x, t] + putdown[x, t]) \leq 1$$

Minimiere die Anzahl der ausgeführten Aktionen:

$$\sum_x \sum_y \sum_t (stack[x, y, t] + unstack[x, y, t]) \\ + \sum_x \sum_t (pickup[x, t] + putdown[x, t]) \longrightarrow \min$$

Übung: Implementieren Sie das Integer-Programm zur Lösung des Brett-Solitär-Problems mittels GLPK und vergleichen Sie dessen Laufzeit mit der Laufzeit eines Programms in C++ für das Problem.

Übung: Erweitern Sie das Integer-Programm zur Lösung des Blocks-World-Problems, so dass mehrere Greifarme gleichzeitig Aktionen durchführen können. Implementieren Sie beide Programme mittels GLPK und vergleichen Sie deren Laufzeiten sowie die berechneten Lösungen.

Wie ist das Integer-Programm zu ändern, wenn der Tisch nur eine begrenzte Anzahl von Blöcken aufnehmen kann?

- *Semantische Netze*
- Regeln
- unsicheres Wissen
- unscharfes Wissen

Domänenwissen: Wissen über

- Objekte und
- ihre Beziehungen.

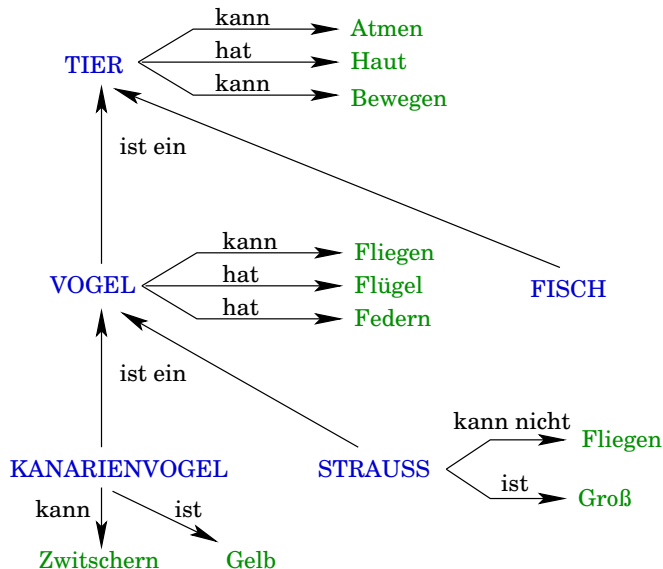
Semantische Netze:

- Stelle Wissen in einem Graphen dar:
 - Knoten entsprechen Fakten bzw. Konzepten und
 - Kanten entsprechen Relationen bzw. Assoziationen zwischen Konzepten.
- Sowohl Knoten als auch Kanten sind in der Regel mit Beschriftungen versehen.

Beziehungen zwischen Objekten können teilweise automatisch erkannt und für die Verarbeitung genutzt werden.

¹⁵aus George F. Luger: Künstliche Intelligenz

Semantische Netze



Semantische Netze entsprechen vielleicht der Speicherung von Informationen beim Menschen. (Harmon und King, 1985)

Satz	Antwortzeit (s)
Ein Kanarienvogel ist ein Kanarienvogel.	1.0
Ein Kanarienvogel ist ein Vogel.	1.18
Ein Kanarienvogel ist ein Tier.	1.26
Ein Kanarienvogel kann zwitschern.	1.32
Ein Kanarienvogel kann fliegen.	1.38
Ein Kanarienvogel hat Haut.	1.48

Am schnellsten konnten die Informationen abgerufen werden, die spezifisch für den jeweiligen Vogel sind.

Auch auf spezifischer Ebene: Verarbeitung von Ausnahmen. Auf die Frage, ob Strauße fliegen können, wurde schneller geantwortet als auf die Frage, ob Strauße atmen können.

Relationen, also Kanten im Netz:

- Hierarchische Relationen:
 - Vererbung
 - Beispiel (Instanz)
 - Teile-Ganzes-Beziehung
- Synonymie (Bedeutungsgleichheit von Ausdrücken)
- Antonymie (Bedeutungsgegensatz von Ausdrücken)
- Kausation: ein Ereignis verursacht ein anderes
- Eigenschaft

Aktuelle Systeme:

- Shapiro: Semantic Network Processing System (SNePS)
- Helbig: MultiNet¹⁶, mehrschichtiges erweitertes semantisches Netz zur Repräsentation natürlichsprachlichen Wissens.

¹⁶http://pi7.fernuni-hagen.de/forschung/multinet/multinet_de.html

Semantic Web: WWW als semantisches Netz

- bisher: WWW verknüpft Daten miteinander → Volltextsuche
- neu: Semantic Web verknüpft Informationen anhand ihrer Bedeutung miteinander → Ontologien

Beispiel:

- Page 1 enthält: Krefeld (Stadt) liegt am Rhein (Fluss).
 - Page 2 enthält: Am Rhein (Fluss) ist es schön (Attribut).
- ⇒ In Krefeld (Stadt) ist es schön (Attribut).
- Dieser Schluss ist möglich, obwohl Informationen in verschiedenen Web-Pages liegen.

Um Krefeld als Ort von Krefeld als Nachname oder Unternehmen zu unterscheiden, müssen Daten mit zusätzlichen Informationen, also Metadaten, versehen werden.

Computer sind mit automatischer Beurteilung oft überfordert.

Beispiel: *Manager* und *Handy* als Suchbegriffe

- Ein Geschäftsführer *ist-ein* *Manager*. (Vererbung)
 - Ein Vorstandsvorsitzender *ist-ein* *Manager*. (Vererbung)
 - Walter Zabel *ist-ein* Geschäftsführer. (Instanz)
 - Ein *Handy* *ist-ein* Mobiltelefon (Synonym)
 - Ein *Handy* *ist-ein* mobiles Endgerät (Synonym)
- Dokumente, in denen Walter Zabel und Mobiltelefon vorkommen, werden gefunden, auch wenn die Suchbegriffe *Manager* und *Handy* nicht im Text vorkommen.

Idee des Semantic Web:

- Inhalte des Hypertexts mit Metadaten beschreiben
 - Ressourcenbeschreibung mittels Resource Description Framework (RDF)
 - Modellierungssprache: Ontology Web Language (OWL)
 - Ziel: Ausdrücke mit wohldefinierten und somit maschinell interpretierbaren Bedeutungen versehen
- inhaltsbezogene Informationssuche

Suchtechnologien im Vergleich:

- Volltextsuche
 - Liefert alle Medieninhalte, die einen oder mehrere der Suchbegriffe beinhalten.
 - Erstellen des Index vollautomatisch → geringe Kosten
 - Darstellen der Ergebnisse in Listenform, schneller Zugriff auf ein Dokument.
- Thesaurus
 - Liefert alle Medieninhalte, die einen oder mehrere der Suchbegriffe oder verwandte, über- oder untergeordnete Begriffe beinhalten.
 - Beispiel: <http://www.dr-antonius.de>
- Semantische Netze
 - Orientiert sich im Gegensatz zu Thesaurus nicht nur an verwandten Begriffen, sondern an allen Assoziationen.
 - Zusammenhänge werden erfasst.

Anfrage:

- Kinderfreundliches Hotel mit Bademöglichkeit am Strand in Norddeutschland.

Linguistische Analyse:

- Hotel: Nomen
- kinderfreundlich: Adjektiv
- mit Bademöglichkeit am Strand: adverbiale Bestimmung
- in Norddeutschland: adverbiale Bestimmung

Synonymersetzung:

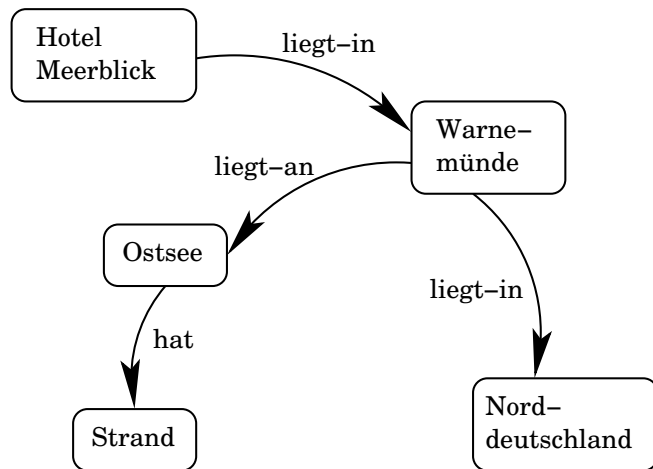
- familienfreundlich entspricht kinderfreundlich

Schlussfolgern:

- Norddeutschland ist-ein Teil von Deutschland
- Deutschland liegt-an Nordsee
- Deutschland liegt-an Ostsee
- Nordsee ist-ein Meer
- Ostsee ist-ein Meer
- Meer hat-ein Strand usw.

Abgleichen:

- Finde Hotels, die der Beschreibung nahe kommen.



→ Ontologie Web Language (OWL)

Ontologie-Schema:

- `:Hotel rdf:type rdfs:Class.`
 - Hotel ist vom Typ „Klasse“, also ein Konzept mit Instanzen
- `:Lokation rdf:type rdfs:Class.`
- `:liegtIn rdf:type owl:TransitiveProperty`
 - „liegtIn“ ist eine Relation (transitiv)
- usw.

Beispiele für Aussagen:

- `:HotelMeerblick rdf:type :Hotel.`
 - „Hotel Meerblick“ ist eine Instanz der Klasse „Hotel“
- `:Warnemünde rdf:type :Lokation.`
 - „Warnemünde“ ist eine Instanz von „Lokation“
- `:Warnemünde :liegtIn :MecklenburgVorpommern.`
- `:MecklenburgVorpommern :liegtIn :Norddeutschland.`
- `:HotelMeerblick :liegtIn :Warnemünde.`
- usw.

Abfragesprache SPARQL:

```
select ?hotel where {?hotel :liegtIn :Warnemünde}
```

- Semantische Netze
- *Regeln*
- unsicheres Wissen
- unscharfes Wissen

Regeln sind formalisierte Konditionalsätze der Form „Wenn A dann B .“ mit der Bedeutung:

Wenn A wahr (erfüllt, bewiesen) ist,
dann schließe, dass auch B wahr ist.

Dabei wird A als *Prämisse* oder *Antezedenz* bezeichnet, B wird *Konklusion* oder *Konsequenz* genannt. Ist die Prämisse einer Regel erfüllt, so sagt man, dass die Regel *feuert*.

Regeln sind eine klassische, weit verbreitete Programmierform für Expertensysteme:

WENN (Batterie ok)
 und (Wert Tankuhr > 0)
 und (Benzinfilter sauber)
DANN (Problem = Zündanlage)

Regeln werden auch gerne in Produktionssystemen zur Steuerung eingesetzt, wobei die Konklusion oft mit einer Aktion verbunden ist:

- *Wenn* der Druck zu hoch ist, *dann* öffne das Ventil.
- *Wenn* die Motortemperatur zu hoch ist, *dann* prüfe die Kühlflüssigkeit.

Solche Regeln werden daher Produktionsregeln genannt.

Modellierung menschlichen Problemlöseverhaltens:

- Erfahrungen basieren auf Problemen, die der Experte früher gelöst hat.
- Diese Erfahrungen hat der Experte in Faustregeln abstrahiert. Wenn ein Experte mit einem neuen Problem konfrontiert wird, wendet er die passende Faustregel an.

Vorteile der Wissensdarstellung in Form von Regeln:

- Regeln stellen einen guten Kompromiss zwischen Verständlichkeit der Wissensdarstellung und formalen Ansprüchen dar.
- Konditionalsätze werden von Menschen seit langer Zeit benutzt, um Handlungsanweisungen oder Prognosen auszudrücken.
- Regeln sind dem Benutzer daher hinreichend vertraut.
- Ein großer Teil von Expertenwissen lässt sich üblicherweise in Regelform ausdrücken.
- Regeln haben eine große Nähe zum menschlichen Denken.

Es gibt verschiedene Arten von Regeln.

- *Produktionsregeln*: Falls in der aktuellen Datenbasis die Bedingung erfüllt ist, führe die Aktion aus.
- *Logische Regeln*: Falls in der aktuellen Datenbasis die Bedingung wahr ist, dann ist auch die Folgerung wahr.

Unterschied zwischen Regeln und IF-THEN-Anweisung in Programmiersprachen:

- Konventionelle Programme haben einen Kontrollfluss: Eine IF-THEN-Anweisung wird nur ausgeführt, wenn das Programm an die entsprechende Stelle kommt.
- In regelbasierten Systemen überprüft die Inferenzkomponente zu jedem Zeitpunkt alle Regeln, ob ihre Bedingung erfüllt ist. Regeln verhalten sich wie WHENEVER-THEN.

Beispiel für eine Regelbasis:

R1	WENN	(Anlasser arbeitet normal)
	DANN	(Batterie ok)
R2	WENN	(Batterie ok)
	und	(Wert Tankuhr > 0)
	und	(Benzinfilter sauber)
	DANN	(Problem = Zündanlage)
R3	WENN	(Batterie ok)
	und	(Wert Tankuhr > 0)
	und	(NICHT Benzinfilter sauber)
	DANN	(Defekt = Benzinzuleitung)
R4	WENN	(NICHT Scheibenwischer ok)
	und	(NICHT Licht ok)
	DANN	(Defekt = Batterie leer)

Fortsetzung:

R5	WENN	(NICHT Wert Tankuhr > 0)
	DANN	(Defekt = Tank leer)
R6	WENN	(Problem = Zündanlage)
	und	(Verteilerdose ok)
	DANN	(Defekt = Zündspule)

Fakten des aktuellen Falls:

F1 Anlasser arbeitet normal

F2 Scheibenwischer ok

F3 Licht ok

F4 Wert Tankuhr > 0

F5 Benzinfilter sauber

F6 Verteilerdose ok

Abgeleitete Fakten:

A1 Batterie ok (R1 mit F1)

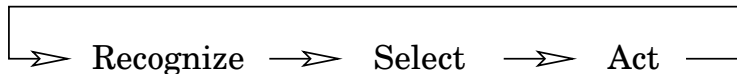
A2 Problem = Zündanlage (R2 mit A1, F4 und F5)

A3 Defekt = Zündspule (R6 mit A2 und F6)

Wenn alle Bedingungen einer Regel durch Fakten der Wissensbasis erfüllt sind, dann führe die Aktion aus.

Recognize-Select-Act Zyklus:

- *Recognize:* Finde alle Regeln, deren Bedingungsteil erfüllt sind.
- *Select:* Wenn mehr als eine Regel anwendbar ist, wird eine Regel ausgewählt (Konfliktlösung).
- *Act:* Führe die Aktion der selektierten Regel aus.



Terminierungsbedingungen: Stop, wenn

- keine Regel mehr anwendbar ist oder
- wenn die Aktion „halt“ ausgeführt wird.

Oft sind wir nicht an allen ableitbaren Fakten interessiert. Daher geht man bei der *Rückwärtsverkettung* von einem Zielobjekt aus, über dessen Zustand der Benutzer Informationen wünscht.

- Das System durchsucht die Regelbasis nach geeigneten Regeln, die das Zielobjekt q in der Konklusion enthalten.
- Die Objekte der Prämisse werden zu Zwischenzielen. Bei einer Regel der Form $p_1 \wedge p_2 \wedge \dots \wedge p_n \rightarrow q$ werden also die Objekte $p_1, p_2, \dots, p_n$ zu Zwischenzielen.
- Die Auswertung der Zwischenziele erfolgt rekursiv bis zu den Fakten. $\rightarrow$ Tiefensuche
- Die Regeln werden also vom Folgerungsteil zum Bedingungsteil hin ausgewertet.
- Dies entspricht der Vorgehensweise eines Prolog-Interpreters.

In der Praxis treten häufig Widersprüchlichkeiten in den Wissensbasen auf:

- Experten benutzen zur Ableitung ihrer Schlüsse oft unausgesprochene Annahmen
- oder übersehen, dass Regeln auch Ausnahmen haben können.

Beispiel:

if V then F	Regel	Alle Vögel können fliegen.
if $V \wedge P$ then $\neg F$	Ausnahme	Pinguine können nicht fliegen.

Wenn in diesem Beispiel sowohl V als auch P wahr sind, erhalten wir zwei Schlüsse: F und $\neg F$

Eine Konsistenzprüfung sollte den Benutzer bereits beim Aufbau einer Regelbasis auf solche Unstimmigkeiten hinweisen.

- Man legt i.d.R. Wert auf eine möglichst einfache syntaktische Form der Regeln.
- Dies ermöglicht eine effizientere Abarbeitung der Regeln und
- es verbessert die Übersichtlichkeit der Wirkung einzelner Regeln.

Üblicherweise stellt man folgende Bedingungen an die Form der Regeln (Prämisse $\Rightarrow$ Konklusion):

- $\vee$ (logisches Oder) darf nicht in der Prämisse einer Regel auftreten.
- Die Konklusion einer Regel sollte nur aus einem Literal bestehen.

Regeln, die diesen Bedingungen nicht genügen, müssen nach den Regeln der klassischen Logik umgeformt werden.

Wenn wir Regeln der Form

$$A \vee B \rightarrow C$$

$$A \rightarrow B \wedge C$$

$$A \rightarrow B \vee C$$

verbieten, dann müssen wir angeben, wie wir solche Regeln umformen:

$$\begin{aligned} A \vee B \rightarrow C &\iff \neg(A \vee B) \vee C && \text{Def. der Implikation} \\ &\iff (\neg A \wedge \neg B) \vee C && \text{de Morgan} \\ &\iff (\neg A \vee C) \wedge (\neg B \vee C) && \text{distributiv} \\ &\iff (A \rightarrow C) \wedge (B \rightarrow C) && \text{Def. der Implikation} \end{aligned}$$

Ebenso kann man herleiten:

$$\begin{aligned} A \rightarrow B \wedge C &\iff \neg A \vee (B \wedge C) && \text{Def. der Implikation} \\ &\iff (\neg A \vee B) \wedge (\neg A \vee C) && \text{distributiv} \\ &\iff (A \rightarrow B) \wedge (A \rightarrow C) && \text{Def. der Implikation} \end{aligned}$$

Nun müssen wir noch herleiten:

$$\begin{aligned} A \rightarrow B \vee C &\iff \neg A \vee (B \vee C) && \text{Def. der Implikation} \\ &\iff (\neg A \vee B) \vee C && \text{assoziativ} \\ &\iff \neg(A \wedge \neg B) \vee C && \text{de Morgan} \\ &\iff A \wedge \neg B \rightarrow C && \text{Def. der Implikation} \end{aligned}$$

Fassen wir oben bei der Anwendung des Assoziativgesetzes anders zusammen, nämlich $(\neg A \vee C) \vee B$, dann erhalten wir:

$$A \rightarrow B \vee C \iff A \wedge \neg C \rightarrow B$$

Fassen wir nun obige Umformungen zusammen und wenden sie auf ein Beispiel an.

Aus der Regel

$$A \vee B \rightarrow C \vee D$$

erhalten wir dann die vier Regeln:

$$A \wedge \neg C \rightarrow D \quad A \wedge \neg D \rightarrow C \quad B \wedge \neg C \rightarrow D \quad B \wedge \neg D \rightarrow C$$

Das ist auch intuitiv klar:

- Die linke Seite ist erfüllt, wenn A oder B wahr ist.
- In diesem Fall muss auch die rechte Seite wahr sein.
- Ist ein Teil der rechten Seite nicht wahr, dann muss der andere rechte Teil erfüllt sein.

Schauen wir uns das an einem konkreten Beispiel aus dem Alltag an.

Beispiel: Wenn es morgen regnet oder schneit, gehen wir ins Kino oder bleiben zu Hause.

Diese Regel lässt sich äquivalent durch die folgenden vier Regeln ausdrücken:

- Wenn es morgen regnet und wir nicht ins Kino gehen, dann bleiben wir zu Hause.
- Wenn es morgen regnet und wir nicht zu Hause bleiben, dann gehen wir ins Kino.
- Wenn es morgen schneit und wir nicht ins Kino gehen, dann bleiben wir zu Hause.
- Wenn es morgen schneit und wir nicht zu Hause bleiben, dann gehen wir ins Kino.

Nun erfüllt jede Regel die gestellten Anforderungen. Die Zahl der Regeln hat sich allerdings erhöht. Der Sinn dieser Umformungen wird im Kapitel Prädikatenlogik erklärt.

Deklarative Programmierung:

- Direkte Repräsentation von Expertenwissen.

Kein expliziter Kontrollfluss:

- Reihenfolge der Regeln spielt keine Rolle.
- Regeln sind aktiv sobald ihre Bedingung erfüllt ist.
- Konfliktlösung wählt eine der aktiven Regeln aus.

Problem:

- Wartung großer Regelsysteme ist aufwändig.
- Auswirkungen von Regeländerungen sind nicht direkt erkennbar.

Regeln sind nicht für alle Arten von Wissen geeignet.

- Semantische Netze
- Regeln
- *unsicheres Wissen*
- unscharfes Wissen

Leider sind Fakten und Regeln nicht immer so eindeutig, wie wir sie bisher dargestellt haben:

- unsichere Fakten:
 - Der Temperatursensor hat eine Toleranz von $\pm 1^{\circ}\text{C}$.
 - Morgen wird es regnen.
- unscharfe Fakten:
 - Ab welcher Temperatur hat man hohes Fieber?
 - Ab wann ist der Druck im Kessel zu hoch?
- unvollständige Fakten: Fehlende Informationen werden durch Annahmen ersetzt.
- unsichere Regeln: „Wenn der Patient Fieber hat, dann hat er sehr wahrscheinlich eine bakterielle Infektion.“

Alle Schlussfolgerungen, die auf solchen Fakten oder Regeln basieren, sind unsicher.

Die klassische zweiwertige Logik mit den Wahrheitswerten 0 und 1 kann hier nicht angewendet werden.

Stattdessen drückt man durch Wahrscheinlichkeiten aus, mit welchem Maß man davon ausgehen kann, dass ein Ereignis eintritt.

Wahrscheinlichkeitsfunktion: Sei Ω ein endlicher Ereignisraum. Eine Funktion $p : 2^\Omega \rightarrow [0, 1]$ heißt Wahrscheinlichkeitsfunktion (oder Wahrscheinlichkeit), wenn gilt:

- $p(\Omega) = 1$
- $p(X \cup Y) = p(X) + p(Y)$ für $X, Y \subseteq \Omega$ und $X \cap Y = \emptyset$

¹⁷aus U. Lämmel, J. Cleve: Lehr- und Übungsbuch Künstliche Intelligenz.

Für solche Wahrscheinlichkeitsfunktionen gilt mit $X, Y \subseteq \Omega$:

- *unmögliches Ereignis*: $p(\emptyset) = 0$
- *Monotonie*: $X \subseteq Y \Rightarrow p(X) \leq p(Y)$
- *Subtraktivität*: für $X \subseteq Y$ gilt:
$$p(Y - X) = p(Y \cap X^c) = p(Y) - p(X)$$
denn $X \subseteq Y \Rightarrow Y = X \cup (Y \cap X^c)$, wobei X und $Y \cap X^c$ disjunkt sind, also gilt:
$$P(Y) = P(X) + P(Y \cap X^c) \iff P(Y \cap X^c) = P(Y) - P(X)$$
- *Subadditivität*: $p(X \cup Y) \leq p(X) + p(Y)$
- *Additivität*: $p(X \cup Y) = p(X) + p(Y) - p(X \cap Y)$
- *Komplement*: $p(X^c) = 1 - p(X)$

Bedingte Wahrscheinlichkeit von X unter der Bedingung B :

$$p(X \mid B) = \frac{p(X \cap B)}{p(B)}$$

Beispiel: Welchen Wahrscheinlichkeitswert muss die Aussage „Wenn eine gerade Zahl gewürfelt wurde, dann ist es eine 2.“ bekommen?

- $X = \{2\}$: Es wurde eine 2 gewürfelt.
- $B = \{2, 4, 6\}$: Es wurde eine gerade Zahl gewürfelt.
- $p(X) = p(\{2\}) = \frac{1}{6}$
- $p(B) = p(\{2, 4, 6\}) = \frac{3}{6} = 0,5$
- $p(X \mid B) = p(\{2\} \mid \{2, 4, 6\}) = \frac{p(\{2\})}{p(\{2, 4, 6\})} = \frac{1}{3}$

Bayessche Formel:

$$p(X | B) = \frac{p(X \cap B)}{p(B)} = \frac{\frac{p(B \cap X)}{p(X)} \cdot p(X)}{p(B)} = \frac{p(B | X) \cdot p(X)}{p(B)}$$

Wenn die Berechnung von $p(X | B)$ verlangt ist, und $p(B | X)$ bekannt ist, dann kann die Schlussfolgerung „umgedreht“ werden.

Beispiel:

- X : Der Motor macht seltsame Geräusche. $p(X) = 0.08$
- Y : Der Motor braucht einen Ölwechsel. $p(Y) = 0.06$
- Ein dringend erforderlicher Ölwechsel sorgt in 60% der Fälle für seltsame Motorengeräusche. $p(X|Y) = 0.6$
- Mit welcher W.-keit deutet ein seltsames Motorgeräusch auf einen dringend erforderlichen Ölwechsel hin?

$$p(Y | X) = \frac{p(X | Y) \cdot p(Y)}{p(X)} = \frac{0.6 \cdot 0.06}{0.08} = 0.45$$

Anmerkung: In praktischen Anwendungen werden oft die Wahrscheinlichkeitstheoretischen Voraussetzungen nicht erfüllt.

- Analytische Wahrscheinlichkeiten oder auf gesicherten statistischen Erhebungen basierende Wahrscheinlichkeiten sind problemlos.
- Rein subjektive Einschätzungen von Experten führen schnell zu Inkonsistenzen.

Wären im obigen KFZ-Beispiel die Wahrscheinlichkeiten subjektiv ermittelt, und nimmt man $p(X) = 0.03$ an, dann ergäbe sich $p(Y|X) = 1.2$ als Wert der bedingten Wahrscheinlichkeit.

Bei subjektiven Einschätzungen ist der Begriff *Glaube* eher angebracht als der Begriff Wahrscheinlichkeit.

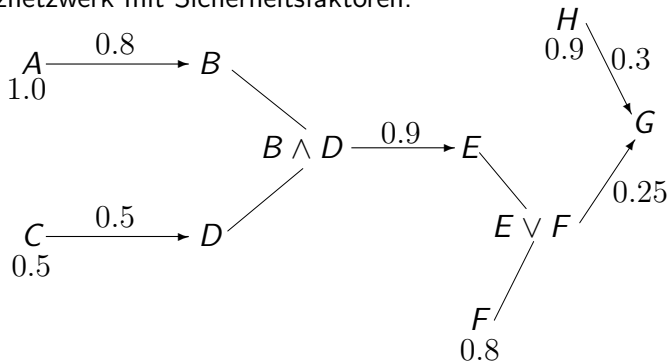
Bereits bei einem der ersten wissensbasierten Systeme überhaupt, dem medizinischen Diagnosesystem MYCIN, musste das Problem der Repräsentation und der Verarbeitung unsicheren Wissens gelöst werden.

- Regeln wurden mit einem Sicherheitsfaktor (certainty factor) versehen, einer reellen Zahl zwischen -1 und $+1$.

Sicherheitsfaktoren drücken den Grad des Glaubens an eine Hypothese aus. $CF(A \rightarrow B)$ gibt an, wie sicher die Konklusion B ist, wenn die Prämisse A wahr ist.

- $CF(A \rightarrow B) = \pm 1$: B ist definitiv wahr/falsch
 - $CF(A \rightarrow B) = 0$: indifferente Haltung
- Auch die Sicherheit eines Fakts F wird mit einem solchen Sicherheitsfaktor $CF(F)$ quantifiziert.
- $CF(B|A)$ beschreibt den Modus Ponens: Aus $(A \rightarrow B$ und $A)$ folgt B , aber mit welcher Sicherheit gilt B ?

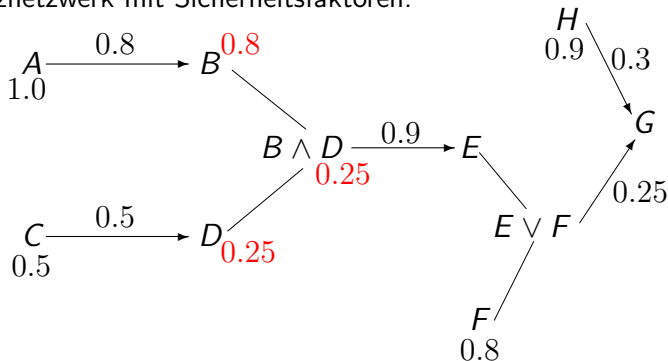
Inferenznetzwerk mit Sicherheitsfaktoren:



In MYCIN wurden folgende Propagationsregeln implementiert:

- Konjunktion: $CF(A \wedge B) = \min\{CF(A), CF(B)\}$
- Disjunktion: $CF(A \vee B) = \max\{CF(A), CF(B)\}$
- Kombination: $CF(B|A) = CF(A \rightarrow B) \cdot \max\{0, CF(A)\}$

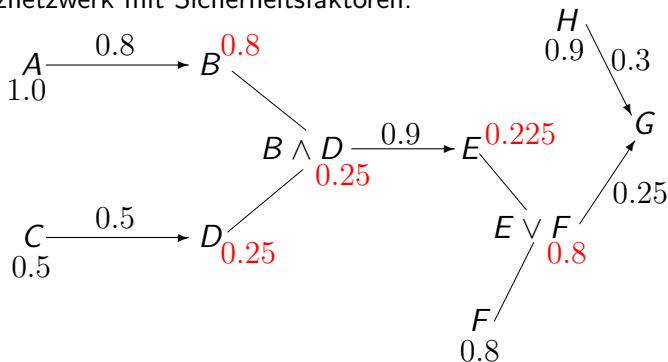
Inferenznetzwerk mit Sicherheitsfaktoren:



Beispiele:

- $CF(B|A) = CF(A \rightarrow B) \cdot CF(A) = 0.8 \cdot 1.0 = 0.8$
- $CF(D|C) = CF(C \rightarrow D) \cdot CF(C) = 0.5 \cdot 0.5 = 0.25$
- $CF(B \wedge D) = \min\{CF(B), CF(D)\} = 0.25$

Inferenznetzwerk mit Sicherheitsfaktoren:



Für G liegen zwei bestärkende Aussagen vor:

- $CF(G|H) = CF(H \rightarrow G) \cdot CH(H) = 0.3 \cdot 0.9 = 0.27$
 - $CF(G|E \vee F) = CF(E \vee F \rightarrow G) \cdot CH(E \vee F) = 0.25 \cdot 0.8 = 0.2$
- ⇒ $CF(G|H, E \vee F)$ ist größer als $CF(G|H)$ und $CF(G|E \vee F)$.

parallele Kombination:

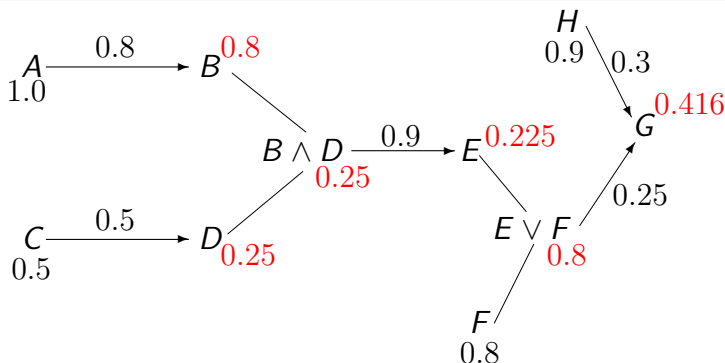
$$\begin{aligned} CF(B|A_1, A_2, \dots, A_n) \\ = f(CF(B|A_1, A_2, \dots, A_{n-1}), CF(B|A_n)) \end{aligned}$$

Dabei ist $f : [0, 1] \times [0, 1] \rightarrow [0, 1]$ definiert als

$$f(x, y) := \begin{cases} x + y - xy & \text{wenn } x, y \geq 0 \\ x + y + xy & \text{wenn } x, y < 0 \\ \text{undefiniert} & \text{wenn } x \cdot y = -1 \\ \frac{x+y}{1-\min\{|x|, |y|\}} & \text{sonst} \end{cases}$$

In unserem Beispiel:

$$\begin{aligned} CF(G|H, E \vee F) \\ = f(CF(G|H), CF(G|E \vee F)) \\ = CF(G|H) + CF(G|E \vee F) - CF(G|H) \cdot CF(G|E \vee F) \\ = 0.27 + 0.2 - 0.054 = 0.416 \end{aligned}$$

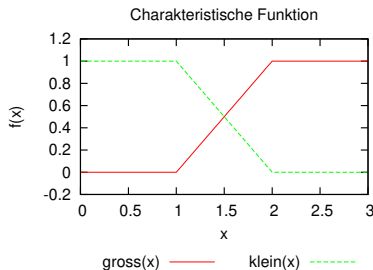


Die Verwendung von Sicherheitsfaktoren ist nur eine heuristische Methode zur Repräsentation und Verarbeitung quantifizierter Unsicherheiten, die sich theoretisch nicht fundieren ließ. Dennoch erwies sie sich als recht leistungsstark und erfolgreich.

- Semantische Netze
- Regeln
- unsicheres Wissen
- *unscharfes Wissen*

Wie können unscharfe Aussagen wie groß/klein, jung/alt oder kalt/heiß modelliert werden?

Charakteristische Funktion: Beschreibt die Zugehörigkeit zu einer Klasse und lässt Werte zwischen 0 und 1 zu.



- Eine 2,10 m große Person, wird von allen als groß angesehen.
- Ein Säugling von 80 cm wird von allen als klein angesehen.
- Ein Volleyballspieler von 2,10 m wird eine 1,80 m große Person nicht als groß bezeichnen, für Kinder wird eine 1,80 m große Person sehr wohl groß sein.

Eine *Fuzzy-Menge* Z über einer Referenzmenge X ist eine Teilmenge aus $X \times [0, 1]$. Dabei wird Z durch die Zugehörigkeitsfunktion $\mu_Z : X \rightarrow [0, 1]$ beschrieben, wobei gilt:

$$z = (x, s) \in Z \iff \mu_Z(x) = s$$

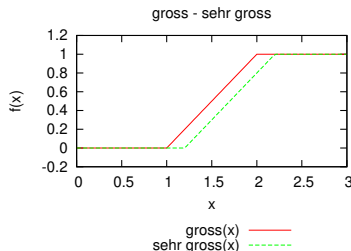
Das Fuzzy-Konzept ist nur eine Verallgemeinerung der üblichen Mengenlehre, bei der wir auch Werte zwischen 0 und 1 erlauben. Klassisch:

$$\mu_A : X \rightarrow \{0, 1\} \quad \text{mit} \quad \mu_A(x) = \begin{cases} 1, & \text{falls } x \in A \\ 0, & \text{sonst} \end{cases}$$

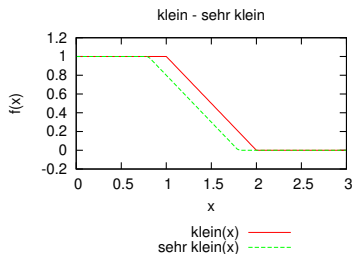
Um die Zugehörigkeit von x zur Fuzzy-Menge *groß* $\cap$ *klein* zur berechnen, müssen die Zugehörigkeitswerte von $\mu_{\text{gross}}(x)$ und $\mu_{\text{klein}}(x)$ geeignet verknüpft werden.

Teilmengenbeziehung: Seien Z_1 und Z_2 zwei Fuzzy-Mengen über X . Z_1 heißt Teilmenge von Z_2 , geschrieben $Z_1 \subset Z_2$, genau dann wenn für alle $x \in X$ gilt:

$$\mu_{Z_1}(x) \leq \mu_{Z_2}(x)$$



sehr groß $\subset$ groß



sehr klein $\subset$ klein

Durchschnitt zweier Fuzzy-Mengen:

Sei $S : [0, 1] \times [0, 1] \rightarrow [0, 1]$ eine Funktion mit den Eigenschaften:

- **neutrales Element:** $S(a, 1) = a$ für $a \in [0, 1]$.

Der Zugehörigkeitswert 1 steht dafür, dass das Element mit Sicherheit zur Fuzzy-Menge Z_2 gehört. Dann ist es plausibel, dass der Zugehörigkeitswert des Durchschnitts dem Zugehörigkeitswert zur Menge Z_1 entspricht.

- **Monotonie:** $a \leq b \Rightarrow S(a, c) \leq S(b, c)$ für $a, b, c \in [0, 1]$.
- **Kommutativität:** $S(a, b) = S(b, a)$ für $a, b \in [0, 1]$.
- **Assoziativität:** $S(a, S(b, c)) = S(S(a, b), c)$ für $a, b, c \in [0, 1]$.

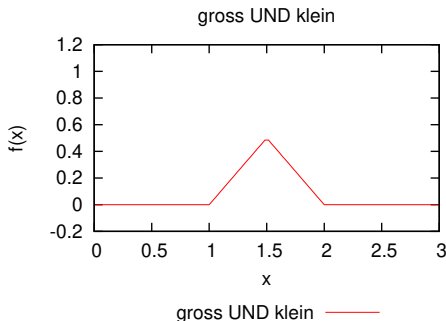
Eine solche Funktion S nennt man *t-Norm*.

Der Durchschnitt zweier Fuzzy-Mengen Z_1 und Z_2 wird definiert durch die Zugehörigkeitsfunktion $\mu_{Z_1 \cap Z_2} = S(\mu_{Z_1}(x), \mu_{Z_2}(x))$.

Mögliche Funktionen:

- $S(a, b) = \min(a, b)$
- $S(a, b) = a \cdot b$

Wählt man das Minimum, erhält man für den Schnitt von *groß* und *klein*:



Vereinigung zweier Fuzzy-Mengen:

Sei $V : [0, 1] \times [0, 1] \rightarrow [0, 1]$ eine Funktion mit den Eigenschaften:

- **neutrales Element:** $V(a, 0) = a$ für $a \in [0, 1]$.

Der Zugehörigkeitswert 0 steht dafür, dass das Element mit Sicherheit nicht zur Fuzzy-Menge Z_2 gehört. Dann ist es plausibel, dass der Zugehörigkeitswert der Vereinigung dem Zugehörigkeitswert zur Menge Z_1 entspricht.

- **Monotonie:** $a \leq b \Rightarrow V(a, c) \leq V(b, c)$ für $a, b, c \in [0, 1]$.
- **Kommutativität:** $V(a, b) = V(b, a)$ für $a, b \in [0, 1]$.
- **Assoziativität:** $V(a, V(b, c)) = V(V(a, b), c)$ für $a, b, c \in [0, 1]$.

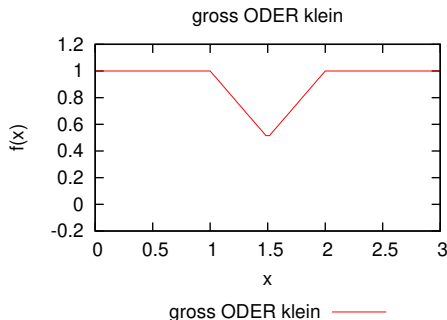
Eine solche Funktion V nennt man *t-CoNorm*.

Die Vereinigung zweier Fuzzy-Mengen Z_1 und Z_2 wird definiert durch die Zugehörigkeitsfunktion $\mu_{Z_1 \cup Z_2} = V(\mu_{Z_1}(x), \mu_{Z_2}(x))$.

Mögliche Funktionen:

- $V(a, b) = \max(a, b)$
- $V(a, b) = a + b - a \cdot b$

Wählt man das Maximum, erhält man für die Vereinigung von *groß* und *klein*:



Die gewählte t-Norm und die t-CoNorm müssen zueinander passen, um die in der klassischen Logik geltenden Aussagen auch bei der Fuzzy-Logik zu erfüllen.

Beispiel: Minimum als t-Norm, Maximum als t-CoNorm. Dann gilt

$$A \vee (A \wedge B) = A$$

auch in der Fuzzy-Logik, denn es gilt:

$$\mu_{A \vee (A \wedge B)}(x) = \max \{ \mu_A(x), \min \{ \mu_A(x), \mu_B(x) \} \} = \mu_A(x)$$

Wie wir aus der Aussagenlogik wissen, sind die Operatoren $\vee$, $\wedge$ und $\neg$ eine Basis. Daher können wir die Implikation $A \rightarrow B$ auch darstellen als $\neg A \vee B$. Die so definierte Zugehörigkeitsfunktion

$$\mu_{A \rightarrow B}(x) := \max\{1 - \mu_A(x), \mu_B(x)\}$$

geht auf Kleene und Dienes zurück. Es gibt weitere Definitionen für die Implikation:

Name	$\mu_{A \rightarrow B}(x)$
Lukasiewicz	$\min\{1, 1 - \mu_A(x) + \mu_B(x)\}$
Kleene-Dienes	$\max\{1 - \mu_A(x), \mu_B(x)\}$
Zadeh	$\max\{1 - \mu_A(x), \min\{\mu_A(x), \mu_B(x)\}\}$
Reichenbach	$1 - \mu_A(x) + \mu_A(x) \cdot \mu_B(x)$

Lotfi A. Zadeh gilt als Erfinder der Fuzzy-Logik (1965), University of California, Berkeley.

Machen wir uns zunächst klar, dass die Implikation $A \rightarrow B$ in der klassischen Aussagenlogik auf zwei verschiedene Arten vorkommt:

- Sind die Wahrheitswerte von A und B bekannt, dann kann auf den Wahrheitswert von $A \rightarrow B$ geschlossen werden.
- Modus Ponens: Wenn bekannt ist, dass die Regel $A \rightarrow B$ gilt, und wir den Wahrheitswert von A kennen, dann können wir auf den Wahrheitswert von B schließen.

In der Fuzzy-Logik ist es genauso: Sind die Zugehörigkeitswerte $\mu_A(x)$ und $\mu_B(x)$ von x bekannt, dann können wir nach einer der obigen Formeln den Zugehörigkeitswert $\mu_{A \rightarrow B}(x)$ von x berechnen.

Beim Schließen mittels Modus Ponens setzen wir die Gültigkeit der Regel voraus, und ermitteln aus dem Zugehörigkeitswert $\mu_A(x)$ den Zugehörigkeitswert $\mu_B(x)$.

Beispiel: Es gelte die Regel „große Menschen sind schwer“, und außerdem gelte „Hans ist groß“. Klassisch:

große Menschen sind schwer	vorhandenes	$A \rightarrow B$
Hans ist groß	Wissen	A
Hans ist schwer	Schlussfolgerung	B

In der Fuzzy-Logik sind unscharfe Werte gegeben. Wenn Hans 1,80 Meter groß ist, dann lesen wir den Zugehörigkeitswert $\mu_{gross}(1,8) = 0,8$ aus unserer Grafik ab: Hans ist ziemlich groß.

große Menschen sind schwer	vorhandenes
Hans ist <i>ziemlich</i> groß	Wissen
Hans ist <i>ziemlich</i> schwer	Schlussfolgerung

Wir übertragen die Werte der Fakten auf die Werte der Folgerung.

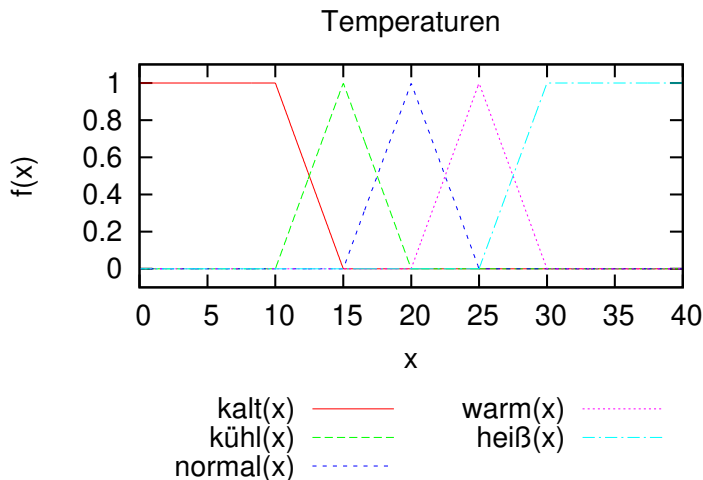
Werte einer *Linguistischen Variablen* sind Worte einer natürlichen Sprache. Bei der Temperatur unterscheiden wir oft die linguistischen Werte *kalt*, *kühl*, *normal*, *warm* und *heiß*.

Diese Werte werden durch unscharfe Mengen A_i bzw. deren Zugehörigkeitsfunktionen $\mu_{A_i}(x)$ repräsentiert.

Die Zugehörigkeitsfunktionen bilden eine linguistische auf eine numerische Werteskala ab.

In der Praxis hat es sich als zweckmäßig erwiesen, die Zugehörigkeitsfunktionen durch stückweisen linearen Verlauf zu definieren. Es sind auch andere Verläufe möglich, wichtig ist nur, dass das Problem adäquat modelliert wird.

Wir weisen den Werten der linguistischen Variablen *kalt*, *kühl*, *normal*, *warm* und *heiß* Fuzzy-Mengen zu:



Wollen wir nun eine Heizungssteuerung implementieren, benötigen wir einige Regeln der Form:

- *Wenn* $T_{innen} = \text{kühl}$ *Und* $T_{ausen} = \text{kalt}$
Dann Heizung = stark
- *Wenn* $T_{innen} = \text{kühl}$ *Und* $T_{ausen} = \text{kühl}$
Dann Heizung = normal
- *Wenn* $T_{innen} = \text{normal}$ *Und* $T_{ausen} = \text{kalt}$
Dann Heizung = normal
- *Wenn* $T_{innen} = \text{normal}$ *Und* $T_{ausen} = \text{kühl}$
Dann Heizung = schwach

Ebenso wie für die Temperatur legen wir linguistische Werte *aus*, *schwach*, *normal*, *stark* und *voll* für die Leistung der Heizung fest.

Zu einer vollständigen Regelung benötigen wir noch weitere Regeln, aber die angegebenen reichen für unser Beispiel.

Wir gehen davon aus, dass wir innen und außen folgende Temperaturen gemessen haben.

$$T_{innen} = 17 \quad \text{und} \quad T_{ausen} = 11$$

Diese Werte sehen wir als scharfe Werte an, so wie die Größe von Hans im vorigen Beispiel.

Wir bestimmen nun die Zugehörigkeitswerte zu den einzelnen Fuzzy-Mengen:

T	$\mu_{\text{kalt}}(T)$	$\mu_{\text{kühl}}(T)$	$\mu_{\text{normal}}(T)$
innen = 17	0	0,6	0,4
außen = 11	0,8	0,2	0

Da wir Regeln der Form $A \wedge B \rightarrow C$ haben, müssen wir zunächst den Zugehörigkeitswert der linken Seite bestimmen. Dazu nutzen wir die Formel $\mu_{A \wedge B}(x) = \min\{\mu_A(x), \mu_B(x)\}$:

außen / innen	kühl (0,6)	normal (0,4)
kalt (0,8)	stark = 0,6	normal = 0,4
kühl (0,2)	normal = 0,2	schwach = 0,2

Wir haben bereits gesehen, dass wir den Zugehörigkeitsgrad der Prämisse für die Schlussfolgerung übernehmen. Welche Werte der linguistischen Variablen *Heizung* in der Tabelle stehen, wird durch die Regeln festgelegt:

$(T_{innen} = \text{kühl} \text{ Und } T_{außen} = \text{kühl}) \longrightarrow \text{Heizung} = \text{normal}$

Problem: Mehrere Regeln führen zu der gleichen Aussage, aber mit unterschiedlichen Zugehörigkeitsgraden.

Wir benötigen eine Verknüpfungsstrategie für die unscharfen Mengen, denen die Regeln in der Konklusion unterschiedliche Zugehörigkeitsgrade zuweisen. Es wurden verschiedene Strategien untersucht:

- Wähle das Maximum der Zugehörigkeitsgrade.
- Wähle den arithmetischen Mittelwert.
- Bilde die Summe der unscharfen Mengen:

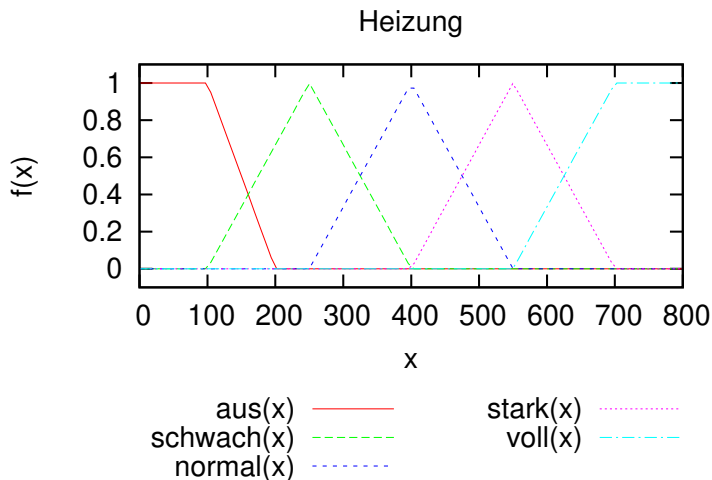
$$\mu_{A+B} = \mu_A + \mu_B - \mu_A \cdot \mu_B$$

Welche Verknüpfungsstrategie gewählt wird hängt davon ab, welche die Prozesscharakteristik am besten beschreibt.

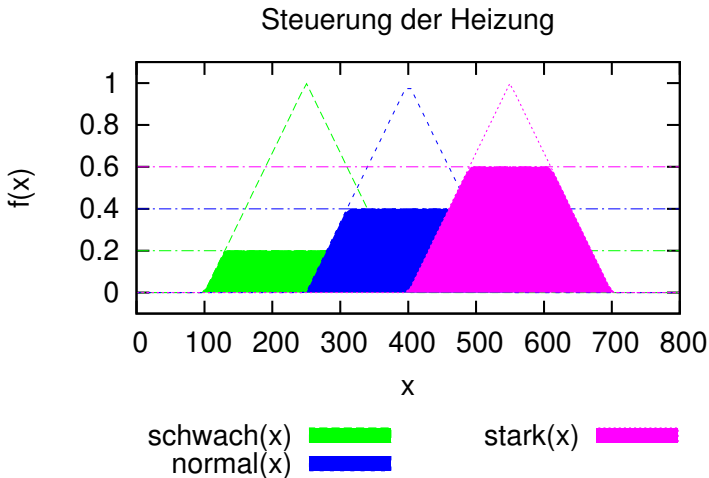
Wir wählen das Maximum der Werte und erhalten so:

$$\mu_{\text{stark}} = 0,6 \quad \mu_{\text{normal}} = 0,4 \quad \mu_{\text{schwach}} = 0,2$$

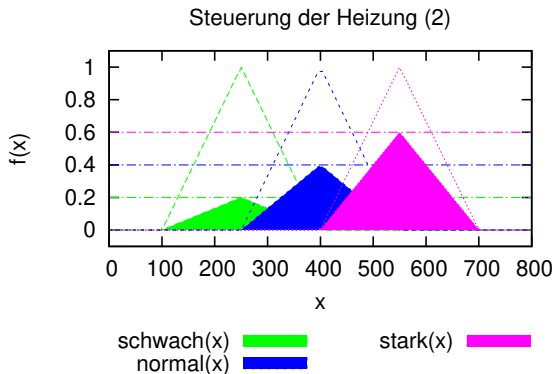
Wenn wir nun die Heizung regeln wollen, müssen wir die Zugehörigkeitswerte de-fuzzifizieren:



Für das De-Fuzzifizieren wurden ebenfalls mehrere Möglichkeiten untersucht. Betrachten wir zunächst die Max/Min-Methode.



Bei der Max/Min-Methode werden die Zugehörigkeitsfunktionen der einzelnen unscharfen Mengen in Höhe des jeweils ermittelten Zugehörigkeitsgrades abgeschnitten, bei der Max/Prod-Methode werden die Werte mit dem Zugehörigkeitsgrad multipliziert. Das bewirkt eine Stauchung der Zugehörigkeitsfunktionen.



Schließlich muss aus diesen Flächen ein einziger Wert ermittelt werden, der dann zur Steuerung genutzt werden kann. Oft wird dafür der Flächenschwerpunkt ermittelt:

$$x_s = \frac{\int_{x_A}^{x_E} x \cdot f(x) dx}{\int_{x_A}^{x_E} f(x) dx}$$

Dabei bedeutet:

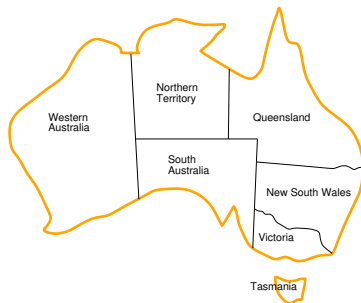
- x_s ist die x -Koordinate des Flächenschwerpunktes
- x_A ist die x -Koordinate des Anfangswertes der Fläche
- x_E ist die x -Koordinate des Endwertes der Fläche
- $f(x)$ ist die Funktion, die die Berandung des Flächenstücks beschreibt

Der so ermittelte Wert x_s wird als Steuergröße für den Regelprozess verwendet.

- Geschichte und Anwendungen
- *Zustandsraumsuche*
- Aussagen- und Prädikatenlogik
- Einführung in PROLOG

- *Motivation*
- Uninformierte Suche
- Informierte Suche
- Lokale Suche
- Game playing

- Sehr viele Probleme der Wissensverarbeitung lassen sich auf ein Suchproblem zurückführen.
- Die Eigenschaften und Lösungsverfahren von Suchproblemen sind daher von grundlegender Bedeutung für die Wissensverarbeitung.
- Suchverfahren sind ein klassisches Kapitel innerhalb der Wissensverarbeitung.
- Wir werden die Zustandsraumsuche im wesentlichen bei Spielen kennen lernen.



Beispiel: Die angegebene Landkarte ist so mit den Farben rot, blau, gelb und grün zu färben, dass keine zwei benachbarten Länder die gleiche Farbe haben.

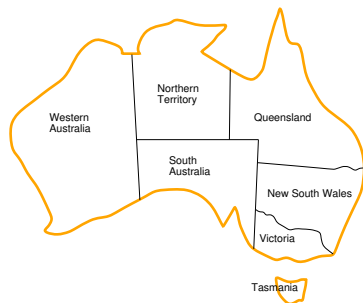
- Ein naives generate-and-test Verfahren würde 4^7 mögliche Farbkonstellationen prüfen.
- Allgemein sind m^n Farbkonstellationen zu prüfen, mit $m :=$ Anzahl der Farben und $n :=$ Anzahl der Länder.

⇒ Ineffizient!

Naives generate-and-test Verfahren:

WA	NT	SA	QL	NSW	Vic	Tas
rot	rot	rot	rot	rot	rot	rot
rot	rot	rot	rot	rot	rot	blau
rot	rot	rot	rot	rot	rot	grün
rot	rot	rot	rot	rot	rot	gelb
rot	rot	rot	rot	rot	blau	rot
rot	rot	rot	rot	rot	blau	blau
rot	rot	rot	rot	rot	blau	grün
rot	rot	rot	rot	rot	blau	gelb
rot	rot	rot	rot	rot	grün	rot
rot	rot	rot	rot	rot	grün	blau
rot	rot	rot	rot	rot	grün	grün
rot	rot	rot	rot	rot	grün	gelb
...	...	...	...	...	...	...

Es scheint sinnvoller zu sein, die Länder der Reihe nach zu färben.

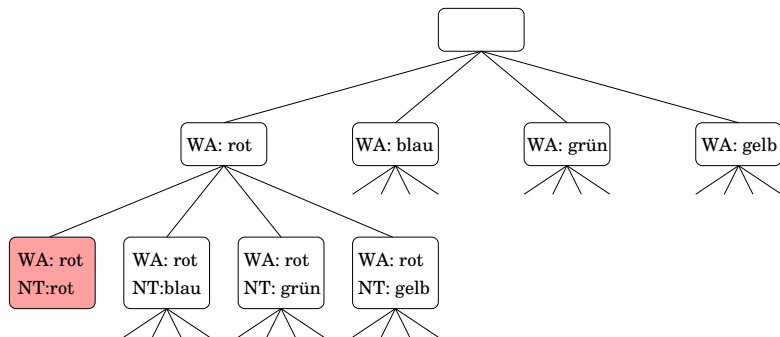


- Beschreibe Zwischenzustände durch Teilfärbungen, etwa (WA:rot, NT:blau, SA:gelb).
- Nach der teilweisen Zuordnung (WA:blau, NT:blau) wird direkt abgebrochen, da die Länder benachbart sind.

- Die Problemlösung startet mit der leeren Färbung ().
- Ziel ist es, eine komplette, zulässige Färbung zu erreichen.
- Die Schritte im Laufe der Problemlösung lassen sich durch Zustandsübergangsoperatoren beschreiben.

Die Lösung des Färbeproblems lässt sich als Suchbaum darstellen:

- Knoten entsprechen den Zuständen (zulässige Teilfärbungen).
- Kanten entsprechen den Operatoren.



- Ungültige Teillösungen werden nicht weiter verfolgt.

Für Suchprobleme lässt sich das Wissen wie folgt repräsentieren:

- Ein Zustand repräsentiert das Wissen zu einem bestimmten Zeitpunkt der Lösungsfindung.
- Die Menge aller Zustände ist der Zustandsraum.
- Zustandsübergangsoperatoren beschreiben den Wechsel von einem Zustand zu einem anderen Zustand des Zustandsraums.
- Den Zustand zu Beginn der Lösungsfindung nennt man Startzustand, er lässt sich explizit angeben.
- Die Menge der Zielzustände charakterisiert die Lösungen des Problems.
- Zielzustände lassen sich in der Regel nur implizit angeben, z.B. über ein Testprädikat.

- Um verschiedene Lösungen vergleichen zu können, können Folgen von Aktionen Kosten zugewiesen werden.

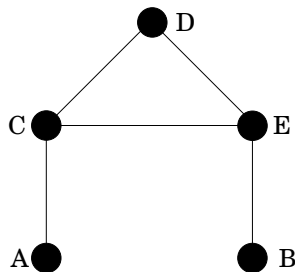
Typischerweise definiert man Kosten abhängig von einem Zustand und einer Aktion.

Die Gesamtkosten für eine Aktionenfolge ergeben sich aus den Einzelkosten.

Routenproblem

Gegeben ist eine Karte mit Städten und Straßen, die die Städte miteinander verbinden.

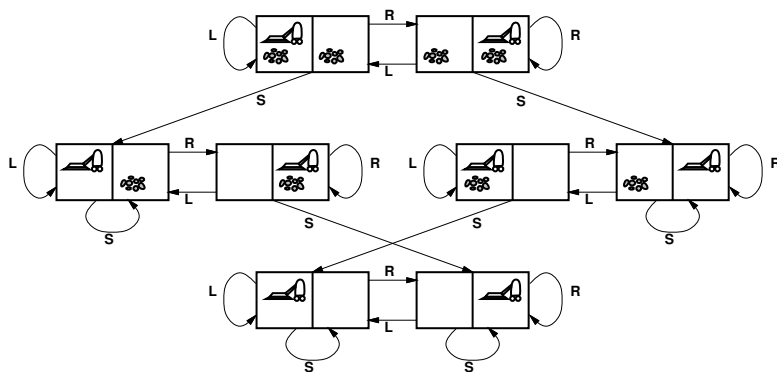
Gesucht ist eine Route von einem Startort zu einem Zielort.



Übung:

- Erstellen Sie den Suchbaum zu obigen Beispiel für den Fall, dass Sie eine Route von A nach B finden wollen.
- Wie können Sie verhindern, dass die Route Kreise enthält?

Staubsaugerwelt



- *Operationen*: rechts, links, saugen
- *Ziel*: alle Räume müssen sauber sein

8-Puzzle

5	4	
6	1	8
7	3	2

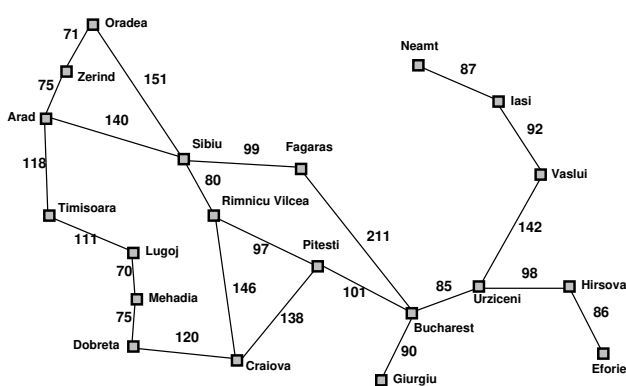
Start State

1	2	3
8		4
7	6	5

Goal State

- *Operationen*: schiebe das leere Feld nach rechts, links, oben oder unten
- *Ziel*: Endkonfiguration

Routenplanung

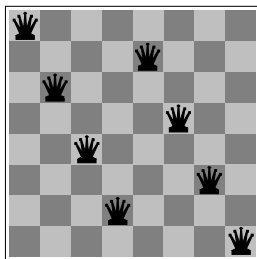


Straight-line distance
to Bucharest

Arad	366
Bucharest	0
Craiova	160
Dobreta	242
Eforie	161
Fagaras	178
Giurgiu	77
Hirsova	151
Iasi	226
Lugoj	244
Mehadia	241
Neamt	234
Oradea	380
Pitesti	98
Rimnicu Vilcea	193
Sibiu	253
Timisoara	329
Urziceni	80
Vaslui	199
Zerind	374

- **Operationen:** Gehe von einer Stadt zu einer benachbarten Stadt entlang einer Kante.
- **Ziel:** Der kürzeste Weg vom Start zum Ziel ist berechnet.
- Der Zielzustand ist hier nicht explizit bekannt. Wie sieht in diesem Fall das Testprädikat aus?

8-Damen-Problem



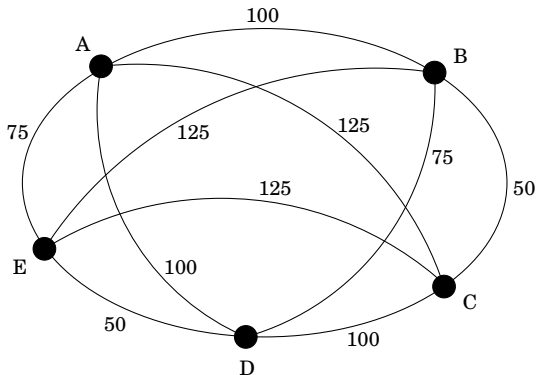
- *Zustand*: Platzierung zwischen 0 und 8 Damen.
- *Operationen*: Eine Dame platzieren.
- *Ziel*: Alle Damen sind so platziert, dass sie sich nicht gegenseitig bedrohen.

Übung:

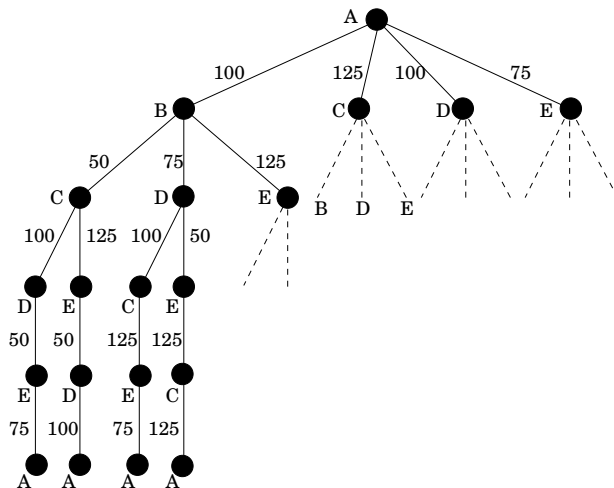
- Auch bei diesem Problem sind die Endzustände nicht explizit bekannt. Wie sieht hier das Testprädikat aus?
- Implementieren Sie eine Zustandsraumsuche und ein naives generate-and-test Verfahren und vergleichen Sie die Lösungen.
- Wie kann der Suchraum verkleinert werden? Bringt das signifikante Laufzeitvorteile?

Problem des Handlungsreisenden

Finde den kürzesten Reiseweg, der genau einmal durch alle Städte führt und wieder bei der ersten Stadt endet.



Problem des Handlungsreisenden



Da die Stadt, in der die Rundreise beginnt, festgelegt ist, hat die erschöpfende Suche bei n Städten eine Komplexität von $(n - 1)!$.

Reduktion der Suchkomplexität durch *Branch-and-Bound*:

- Generiere zu einem Zeitpunkt jeweils nur einen Weg.
- Die Kosten des bisher besten Weges wird als Schranke bei weiteren Kandidaten eingesetzt.
- Ist ein Teilweg bereits teurer als die bisher beste Lösung, so werden alle Erweiterungen des Weges nicht weiter betrachtet.
- Weitere Reduktion durch Schätzen der verbleibenden Kosten:
 - Jede der noch nicht besuchten Städte muss besucht und auch wieder verlassen werden.
 - Addiere jeweils die Werte der zwei kostengünstigsten Kanten, die inzident zu einem noch unbesuchten Knoten sind.
 - Die Summe muss noch durch 2 dividiert werden, da jede Kante zweimal genutzt wird: um einen Knoten zu verlassen und um den nächsten zu besuchen.

→ Komplexität: 1.26^n , leider immer noch exponentiell!

Die Berechnung der Nachfolger eines Knotens s wird als *Expansion* des Knotens s bezeichnet.

- Der Zustandsraum beschreibt lediglich, *wie* man zu einer Lösung gelangen kann, aber nicht, wie man effizient zu dieser kommt.
- Wesentlich für eine effiziente Problemlösung sind:
 - Die Reihenfolge, in der die Zustände untersucht bzw. expandiert werden
 - sowie die Bewertung der einzelnen Zustände.

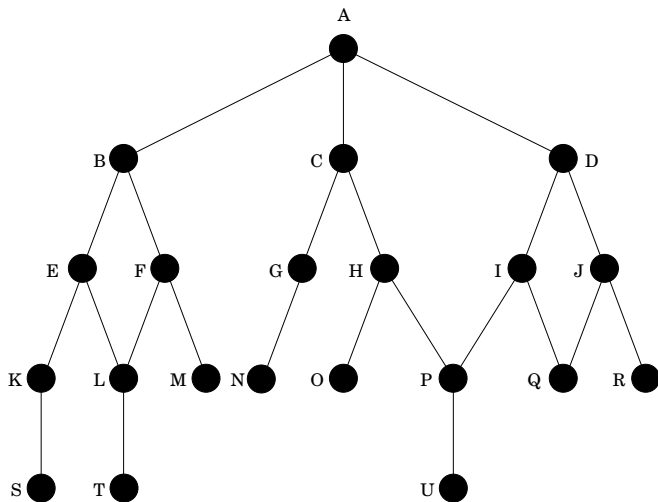
- Motivation
- *Uninformierte Suche*
- Informierte Suche
- Lokale Suche
- Game playing

- Suchverfahren, die über die Beschreibung des Zustandsraums hinaus keine Zusatzinformation benutzen.
- Insbesondere findet keine Bewertung der einzelnen Zustände statt. Um Wiederholungen zu vermeiden, werden bereits untersuchte Zustände markiert oder gespeichert.
- Die Verfahren unterscheiden sich im wesentlichen darin, in welcher Reihenfolge die Zustände expandiert werden.
- Die wichtigsten Vertreter der uninformierten Suchverfahren sind die *Breitensuche* und die *Tiefensuche*.
- Ausgehend von der Wurzel des Suchbaums (Startzustand) werden die Knoten sukzessive expandiert.
- Später wird von den Nachfolgern des expandierten Knotens weiter gearbeitet, solange bis ein Zielknoten gefunden wurde.

Breiten- und Tiefensuche laufen nach dem gleichen Schema ab:

- *Agenda*: Liste der Knoten, die gerade in Bearbeitung sind.
- Knoten der Agenda sind generiert, aber noch nicht expandiert.
- Zu Beginn der Suche besteht die Agenda nur aus dem Startzustand.
- In einer beliebigen Iteration wird der erste Knoten s_{akt} aus der Agenda genommen.
- Wenn s_{akt} ein Zielzustand ist, hat man eine Lösung gefunden.
- Ist s_{akt} kein Zielzustand, so wird s_{akt} expandiert, d.h. alle Nachfolger von s_{akt} werden in die Agenda eingefügt.
- Breiten- und Tiefensuche unterscheiden sich darin, wo die Nachfolger in die Agenda eingefügt werden.

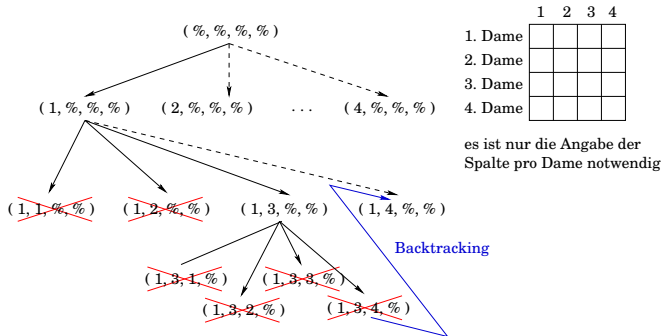
Beispiel für Tiefen- und Breitensuche



- Bei der Tiefensuche werden die Nachfolger eines expandierten Knotens an den Anfang der Agenda eingefügt.
- Die Agenda entspricht einem Kellerspeicher (Stack).

```
agenda = new Stack();
agenda->push(startknoten);
while (!agenda->isEmpty()) {
    akt = agenda->pop();
    if (zielknoten(akt))
        return akt;
    foreach node in nachfolger(akt)
        agenda->push(node);
}
return 0;
```

- Liefert ein Knoten, der kein Zielknoten ist, keine neuen Knoten, so wird die Suche fortgesetzt an dem nächstgelegenen Knoten, für den noch nicht alle Nachfolger expandiert wurden.
- Verhalten entspricht einem Backtracking (zurückverfolgen).

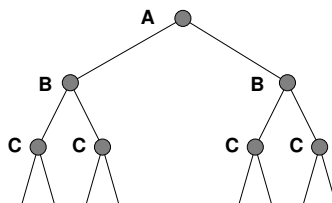
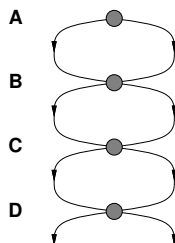


- Die Abarbeitung wird an den Knoten abgebrochen, die zu keiner zulässigen Lösung führen können.

Repeated States

Wie verhindern wir, dass Zustände wiederholt untersucht werden?

- Problem tritt immer dann auf, wenn Operationen umkehrbar sind: Schiebepuzzle, Routenproblem, Spiele, usw.
- Wird nicht erkannt, dass Zustände bereits untersucht wurden, kann aus einem linearen Problem ein exponentielles werden:



- Der einzige Weg, wiederholtes Untersuchen von Zuständen zu vermeiden, ist: Speichere bereits besuchte Zustände.

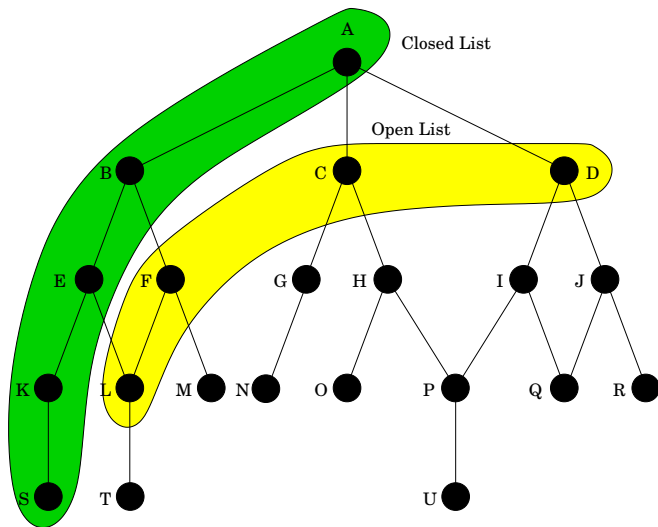
→ Tradeoff between space and time!

- *Algorithms that forget their history are doomed to repeat it.*
- Um die bereits untersuchten Zustände zu speichern, werden expandierte Knoten in der *closed list* verwaltet.
- Wenn der aktuell zu untersuchende Zustand bereits in der closed list gespeichert ist, wird er verworfen anstatt expandiert zu werden.
- Die Agenda wird daher auch *open list* genannt.

Die um eine closed list erweiterte Tiefensuche zur Vermeidung wiederholter Zustände:

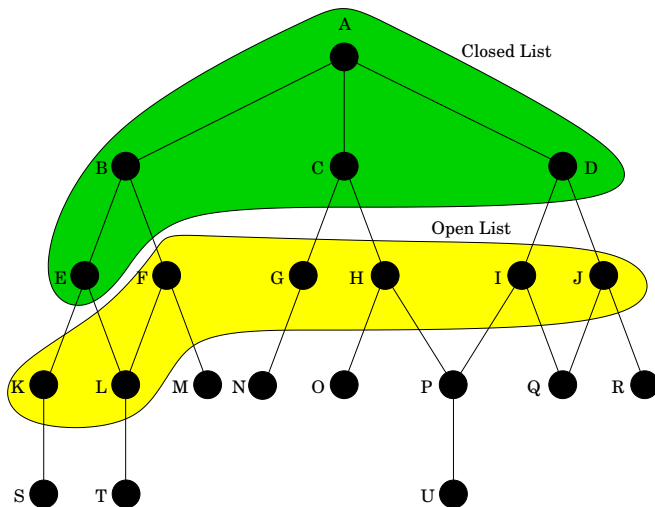
```
agenda = new Stack();
agenda->push(startknoten);
closed = new HashMap();
while (!agenda->isEmpty()) {
    akt = agenda->pop();
    if (zielknoten(akt))
        return akt;
    if (!closed->contains(akt)) {
        closed->insert(akt);
        foreach node in nachfolger(akt)
            agenda->push(node);
    }
}
return 0;
```

Tiefensuche



- Bei der Breitensuche werden die Nachfolger eines expandierten Knotens an das Ende der Agenda eingefügt.
- Die Agenda entspricht einer Warteschlange (Queue).

```
open = new Queue();
open->append(startknoten);
closed = new HashMap();
while (!open->isEmpty()) {
    akt = open->first();
    if (zielknoten(akt))
        return akt;
    if (!closed->contains(akt)) {
        closed->insert(akt);
        foreach node in nachfolger(akt)
            open->append(node);
    }
}
return 0;
```



Ein Weinhändler hat drei Krüge mit 9l, 7l bzw. 4l Inhalt.

- Auf den Krügen sind keine Litermarkierungen angebracht.
- Der 9l-Krug ist gefüllt, die anderen beiden sind leer.
- Die Krüge sollen so umgefüllt werden, dass der 9l-Krug sechs Liter und der 4l-Krug drei Liter enthält.

Erstellen Sie den Suchbaum und ermitteln Sie mittels Tiefen- und Breitensuche eine Lösung des Problems.

Betrachten Sie das Kannibalen und Missionare-Problem:

- 3 Missionare und 3 Kannibalen wollen einen Fluss überqueren.
- Es steht nur ein Ruderboot für die Überfahrt zur Verfügung.
- Das Boot kann maximal 2 Personen aufnehmen.
- Es dürfen nie mehr Kannibalen als Missionare an einem Ort sein, sonst gibt es ein großes Fressen.

Aufgabe:

- Definieren Sie die Zustandsmenge und die Operatoren, die einen Zustand in einen anderen Zustand überführen.
- Erstellen Sie den Suchbaum und ermitteln Sie mittels Tiefen- und Breitensuche eine Lösung des Problems. Zustände, bei denen mehr Kannibalen als Missionare an einem Ort sind, werden nicht weiter verfolgt.

Ungünstiger Fall für Breiten- und Tiefensuche: Die Lösung liegt in der „äußersten, unteren rechten Ecke“ des Suchbaums.

- Sei b die Verzweigungsrate und
- sei t die Tiefe des Zielknotens.

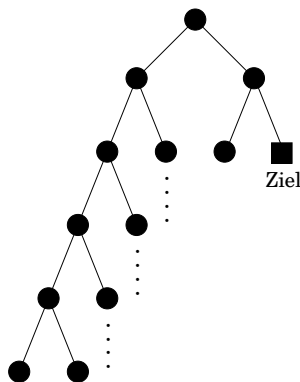
Breitensuche: Die Agenda kann eine komplette Ebene des Suchbaums enthalten, daher erhalten wir:

- Zeitkomplexität: $\Theta(b^t)$
- Platzkomplexität: $O(b^t)$

Tiefensuche: Sei m die maximale Tiefe des Baums.

- Platzkomplexität: $O(b \cdot m)$, denn die Agenda enthält nur die Knoten des aktuellen Suchpfades sowie deren Nachfolger.
- Zeitkomplexität: $O(b^m)$

Laufzeit Tiefsuche: $O(b^m)$



$b = 2$ (Verzweigungsrate)

$t = 2$ (Tiefe des Ziels)

Eine Breitensuche würde das Ziel viel schneller finden. Leider ist bei vielen Problemen keine Breitensuche einsetzbar, z.B. bei Schach oder Go, da dort die Verzweigungsrate zu hoch ist!

Ein Suchverfahren heißt *vollständig*, wenn für jeden Suchbaum jeder Knoten expandiert werden könnte, solange noch kein Zielknoten gefunden wurde.

- Ein vollständiges Suchverfahren ist fair in dem Sinne, dass jeder Knoten die Chance hat, expandiert zu werden.
- Ein vollständiges Suchverfahren findet auch bei unendlichen Suchbäumen stets eine Lösung, falls eine existiert.

Nach dieser Definition gilt:

- Die Breitensuche ist vollständig.
- Tiefensuche ist nur bei endlichen Suchbäumen vollständig.
Hat ein Zweig im Suchbaum unendliche Tiefe, dann können Knoten in anderen Zweigen nicht expandiert werden.

Für ein uninformatiertes Suchverfahren heißt eine Lösung *optimal*, wenn sie unter allen Lösungen die geringste Tiefe im Suchbaum aufweist.

Nach dieser Definition gilt:

- Die Breitensuche findet eine optimale Lösung, falls eine Lösung existiert.
- Eine Tiefensuche findet im Allgemeinen keine optimale Lösung.

Wir versuchen, die Nachteile der Tiefensuche auszuschalten, indem wir eine maximale Suchtiefe vorgeben:

- Zweige jenseits der maximalen Suchtiefe werden abgeschnitten.
- Dadurch wird garantiert, dass eine Lösung gefunden wird, wenn eine innerhalb der Suchtiefe existiert.

Frage: Stimmt das so ohne Weiteres? Was ist mit den in der Closed-List gespeicherten Zuständen?

- Nachteil: Wie groß soll man die maximale Suchtiefe wählen?

Da man nicht weiß, wie groß man die maximale Suchtiefe wählen soll, probiert man dies einfach systematisch durch: schrittweise vertiefende Suche

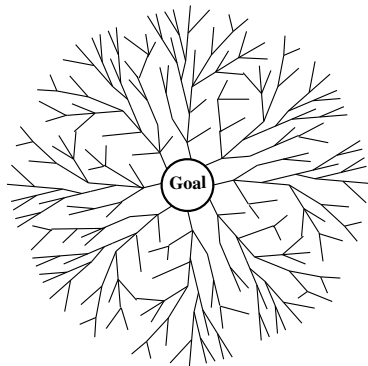
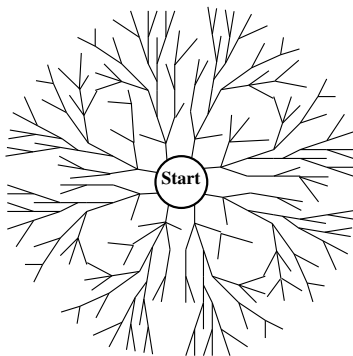
- Führe immer wieder eine Tiefensuche mit Begrenzung der Suchtiefe durch und erhöhe nach einer erfolglosen Suche die maximale Suchtiefe jeweils um eins.
- Bei einer maximalen Suchtiefe t und einem regulären Suchbaum mit Verzweigungsgrad b werden $O(b^t)$ Knoten erzeugt.
- Man verliert bei diesem Verfahren nur einen konstanten Faktor¹⁸:

$$\sum_{d=0}^t b^d = \frac{b^{t+1} - 1}{b - 1} < \frac{b}{b - 1} b^t = O(b^t)$$

¹⁸dies gilt nur, falls der Verzweigungsgrad konstant ist, also die Suche einen exponentiellen Zeitaufwand hat

Bidirektionale Suche

- Gleichzeitige Suche vom Start zum Ziel und vom Ziel zum Start.
- Lösungen ergeben sich dort, wo sich die Suchpfade treffen.



Durch Halbierung der Länge der Suchpfade ergibt sich ein Aufwand¹⁹ von $O(2b^{\frac{d}{2}}) = O(b^{\frac{d}{2}})$.

⇒ Drastische Reduzierung der Laufzeit! Die in akzeptabler Zeit lösbare Problemgröße kann verdoppelt werden.

Probleme:

- Zielzustände sind oft nicht explizit bekannt: n -Damen-Problem
- Zielzustände müssen nicht eindeutig sein: n -Damen-Problem
- Operatoren sind nicht unbedingt umkehrbar.

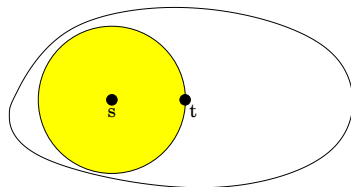
Übung: Implementieren Sie eine bidirektionale Suche für Brett-Solitär. Ergibt sich eine signifikante Verbesserung der Laufzeit gegenüber Ihrem ursprünglichen Programm?

¹⁹dies gilt nur, falls der Verzweigungsgrad konstant ist, also die Suche einen exponentiellen Zeitaufwand hat

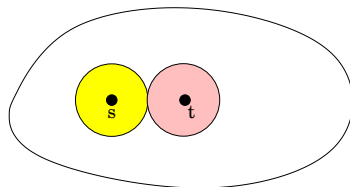
Bidirektionale Suche

Bei der Routenplanung mittels Dijkstras Algorithmus kann die Berechnung durch die bidirektionale Suche nur um den Faktor 2 beschleunigt werden.

bisher: $d(s, t)^2 \cdot \pi$

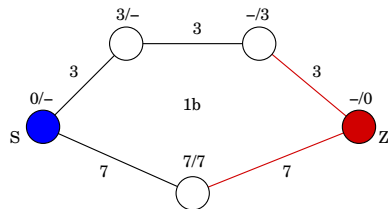
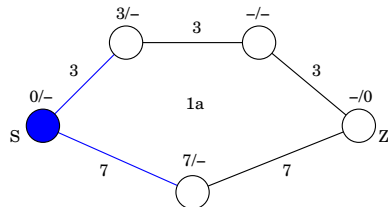
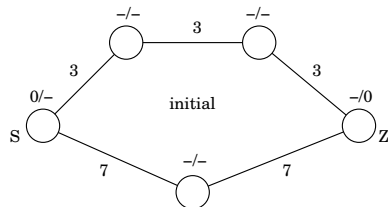


jetzt: $2 \cdot \left(\frac{d(s, t)^2}{4} \cdot \pi \right)$



Bidirektionale Suche

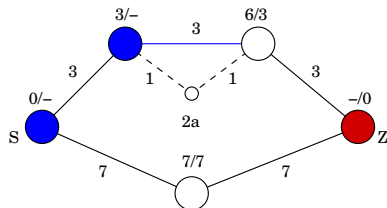
Außerdem darf bei der Routenplanung die Suche nicht dann abgebrochen werden, wenn sich zwei Suchpfade treffen.



Die Suchpfade treffen sich zwar in dem unteren Knoten, aber der kürzeste Weg wurde noch nicht gefunden.

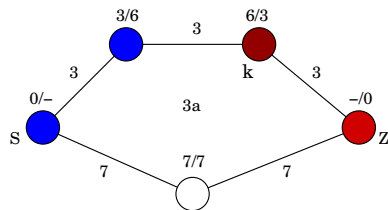
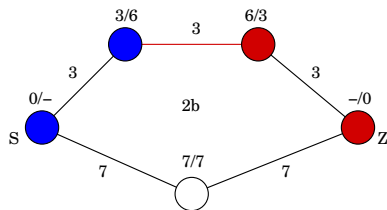
Die Suche wird also fortgesetzt.

Bidirektionale Suche



Auch hier darf die Suche noch nicht abgebrochen werden, denn es könnte noch einen kürzeren Weg geben.

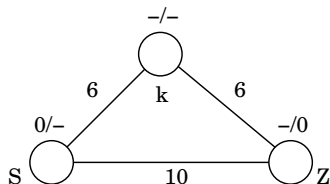
Die Suche wird also fortgesetzt.



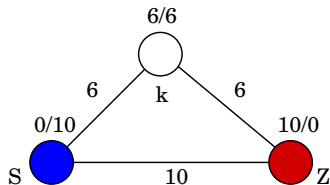
Erst wenn ein Knoten k von beiden Suchrichtungen aus als final markiert wurde, kann die Suche abgebrochen werden.

Frage: Ist k auf kürzesten Weg vom Start zum Ziel enthalten?

Nein! Ein Knoten k , der von beiden Suchrichtungen aus als final markiert wurde, muss nicht notwendig auf einem kürzesten Weg vom Start zum Ziel enthalten sein.



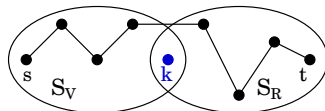
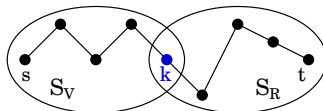
Bei beiden Suchrichtungen wird im ersten Schritt jeweils die Distanz zum Knoten k auf den Wert 6 aktualisiert.



Im nächsten Schritt wird jeweils der Knoten k ausgewählt und aus der Priority-Queue entfernt. Aber k liegt nicht auf dem kürzesten Weg vom Start S zum Ziel Z .

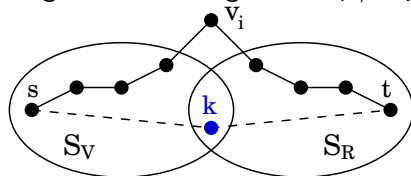
Aber: Der kürzeste Weg wurde korrekt berechnet, sobald ein Knoten k gefunden wurde, der von beiden Suchen als final markiert wurde.

- natürlich gilt: $\overleftarrow{d}[v] + \overrightarrow{d}[v] \geq \text{dist}(s, t)$
- Seien S_V und S_R die abgearbeiteten Knoten der Vorwärts- bzw. Rückwärtssuche nach Terminierung.
- Sei $P = (v_1, \dots, v_m)$ ein kürzester Weg vom Start zum Ziel.
- zeige:
 - $\{v_1, \dots, v_m\} \subseteq S_V \cup S_R$ oder



- $|P| = \text{dist}(s, k) + \text{dist}(k, t)$

- Angenommen, es gibt ein $v_i \notin S_V \cup S_R$



- Dann gilt $dist(s, v_i) \geq dist(s, k)$ und $dist(v_i, t) \geq dist(k, t)$, da sonst Dijkstra Knoten v_i ausgewählt hätte und nicht k .
- Also ist $|P| = dist(s, v_i) + dist(v_i, t) \geq dist(s, k) + dist(k, t)$.
- Da P nach Voraussetzung ein kürzester Weg ist, muss gelten: $|P| = dist(s, v_i) + dist(v_i, t) = dist(s, k) + dist(k, t)$

- Motivation
- Uninformierte Suche
- *Informierte Suche*
- Lokale Suche
- Game playing

- Für größere Suchbäume sind Breiten- und Tiefensuche nicht effizient genug.
- Vielversprechender sind Ansätze, bei denen Problemwissen zur Steuerung des Suchprozesses eingesetzt wird.
- Dies kann dadurch geschehen, dass die Zustände (Knoten) danach bewertet werden, wie erfolgversprechend sie sind.
- Man schätzt bspw. für jeden Knoten, wie nahe er an einem Zielknoten liegt.
- Solch eine Bewertung nennt man heuristische Funktion.
 - griech. eurisco heißt „ich entdecke“
 - „Heureka“: ich habe es gefunden
- Heuristik bei Zustandsraumsuche: Regeln zur Auswahl solcher Zweige, die mit hoher Wahrscheinlichkeit zu einer akzeptablen Problemlösung führen.

Wir lernen die Zustandsraumsuche im wesentlichen bei Spielen kennen, weil

- die Suchräume groß genug sind, um eine heuristische Beschneidung (engl. Pruning) zu erfordern.
- die meisten Spiele komplex genug sind, so dass eine Vielzahl heuristischer Bewertungen für den Vergleich und die Analyse einsetzbar sind.
- Spiele i.Allg. keinen komplexen Repräsentationsaufwand erfordern, so dass man sich auf das Verhalten der Heuristik statt auf Probleme der Wissensrepräsentation konzentrieren kann.
- alle Knoten des Zustandsraums eine gemeinsame Repräsentation besitzen, so dass im gesamten Suchraum eine einzige Heuristik angewendet werden kann.

Wann werden informierte Suchverfahren eingesetzt?

- Wenn ein Problem keine exakte Lösung hat, weil bspw. die Problembeschreibung oder die verfügbaren Daten zweideutig sind.
 - Medizinische Diagnose: Eine Menge von Symptomen kann verschiedene Ursachen haben.
 - Wähle die wahrscheinlichste Lösung.
 - Visuelle Wahrnehmung: Optische Täuschung, Ausrichtung oder Orientierung von Objekten nicht klar.
 - Wähle die wahrscheinlichste Interpretation der Szene.
- Wenn die Rechenkosten für die Suche nach einer exakten Lösung zu groß sind, weil die Anzahl der Zustände exponentiell oder faktoriell mit der Tiefe der Suche wächst, wie z.B. bei Schach oder Go.
 - Suche entlang eines „viel versprechenden“ Weges.

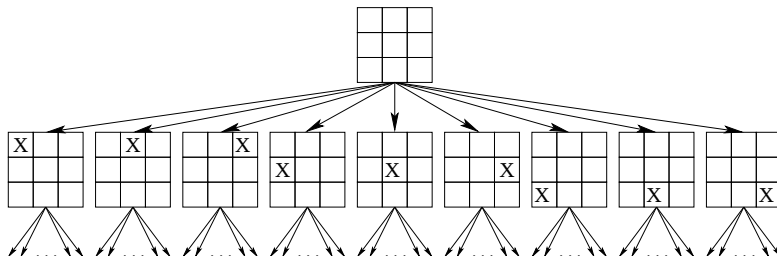
Anwendungen:

- Spiele: Es kann nicht jede erreichbare Spielsituation beim Schach oder Go weiterverfolgt werden.
- Theorembeweise: Es kann nicht jede Schlussfolgerung weiterverfolgt werden.

Expertensysteme:

- Heuristiken sind essentielle Komponenten von Problemlöseverfahren.
 - Menschliche Experten verwenden Faustregeln.
 - Faustregeln sind weitgehend heuristisch.
- Extrahiere und formalisiere Faustregeln.

Tic-Tac-Toe ohne Heuristik:

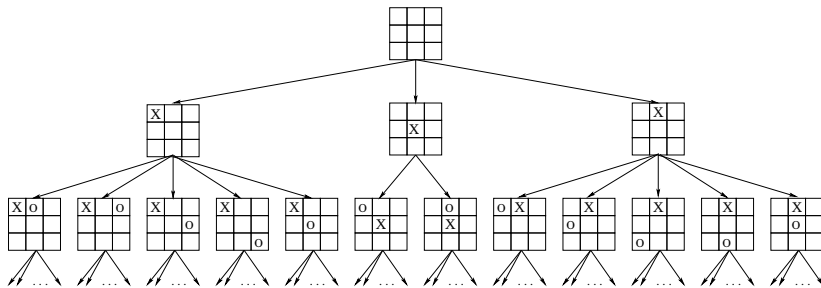


⇒ $9! = 9 \cdot 8 \cdot 7 \cdot \dots \cdot 2 \cdot 1 = 362.880$ mögliche Zustände

Beachte: Viele doppelte Zustände durch Zuggenerierung. Es gibt weniger als $3^9 = 19.683$ verschiedene Zustände!

⇒ Es scheint sinnvoll zu sein, bereits untersuchte Zustände in einer Hash-Tabelle zu verwalten.

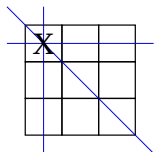
Reduktion bzgl. Symmetrien bei Tic-Tac-Toe:



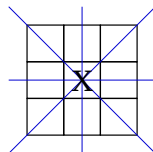
⇒ $12 \cdot 7! = 60.480$ mögliche Zustände

Problem: Wie baut man einen Zuggenerator, der Symmetrien berücksichtigt?

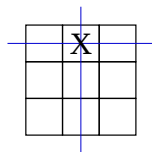
Heuristik der „meisten Gewinnmöglichkeiten“ bei Tic-Tac-Toe:



Anzahl Gewinn-
möglichkeiten 3



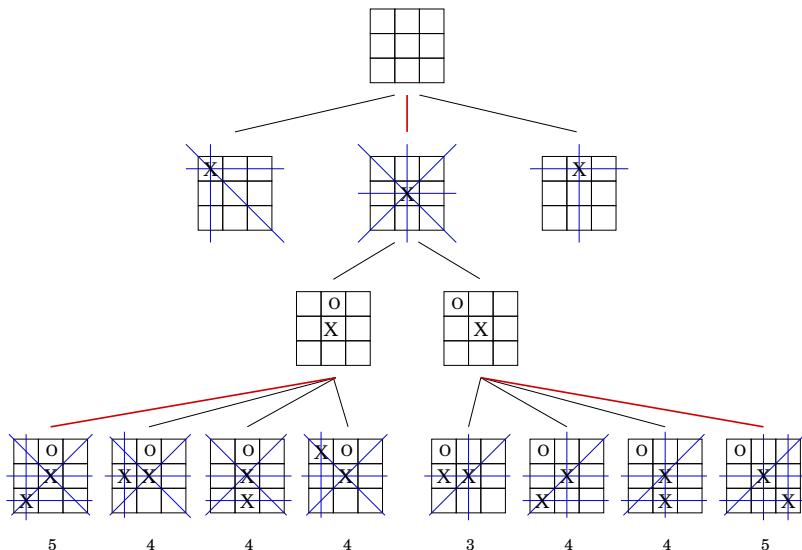
4



2

Wir beschränken also die Suche auf die Nachfolgezustände des obigen, mittleren Zustands. Anschließend wenden wir noch die Reduktion bzgl. Symmetrien an.

Informierte Suchverfahren



⇒ höchstens $9 \cdot 8 = 72$ mögliche Zustände

Heuristische Funktion

Eine Funktion, die jedem Zustand (Knoten) s eines Zustandsraums (Suchbaums) eine nichtnegative Zahl $h(s)$ zuordnet, heißt *heuristische Funktion*. Für einen Zielzustand s gilt dabei $h(s) = 0$.

Ein Suchverfahren, das eine heuristische Funktion zur Auswahl der zu expandierenden Zustände einsetzt, heißt *informiertes Suchverfahren* oder auch *heuristisches Suchverfahren*.

Beispiel Schiebepuzzle:

5	4	
6	1	8
7	3	2

Start State

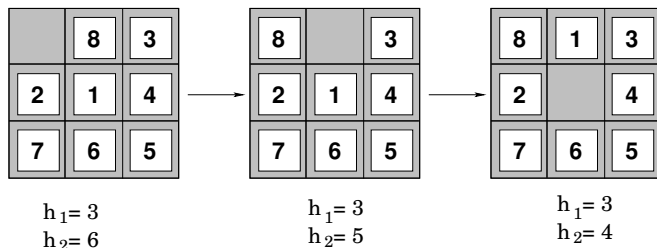
1	2	3
8		4
7	6	5

Goal State

Mögliche heuristische Funktionen:

- $h_1(s)$: Anzahl der Plättchen, die nicht an der richtigen Stelle liegen. Hier: $h_1(s) = 7$
- $h_2(s)$: Summe der Entfernungen aller Plättchen von der Zielposition, also die Manhattan-Distanz. Hier:
 $h_2(s) = 2 + 3 + 3 + 2 + 4 + 2 + 0 + 2 = 18$

Heuristische Funktion

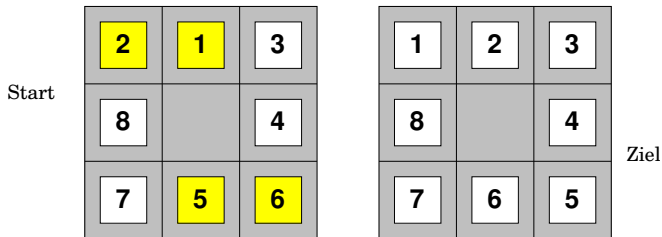


- Die heuristische Funktion h_2 differenziert stärker als h_1 , d.h. h_2 kann Zustände unterscheiden, die von h_1 gleich bewertet werden.
- Eine heuristische Funktion ist um so brauchbarer, je mehr Zustände sie unterschiedlich bewertet.
- Eine heuristische Funktion, die alle Zustände gleich bewertet, ist unbrauchbar.

Heuristische Funktion

Beide Heuristiken sind zu kritisieren, da sie die Schwierigkeiten eines Spielsteintauschs nicht berücksichtigen:

- Es sind wesentlich mehr als zwei Verschiebungen notwendig, da die Plättchen umeinander herum verschoben werden müssen.
- Man kann eine kleine Zahl wie 2 auf die Heuristik multiplizieren, um dies zu berücksichtigen.



- $h_3(s) = 2 \times \text{Anzahl direkter Spielsteintausch}$

- h_3 versagt, wenn ein Zustand keinen direkten Tausch enthält:

2	8	3
1	6	4
	7	5

$h_1: 5$ $h_2: 6$ $h_3: 0$

2	8	3
1		4
7	6	5

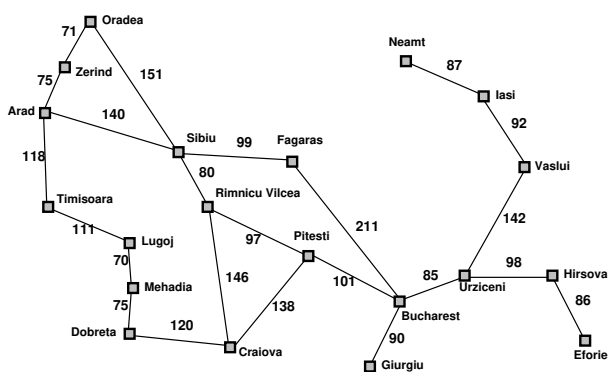
$h_1: 3$ $h_2: 4$ $h_3: 0$

2	8	3
1	6	4
7	5	

$h_1: 5$ $h_2: 6$ $h_3: 0$

- Heuristik h_4 : Die Summe der Entfernungen plus das Doppelte der Anzahl direkter Tauschkonfigurationen.
- Heuristik: Begrenzte Information über Zustandsbeschreibung einsetzen, um intelligente Auswahl zu treffen.
- Jede obige Heuristik ignoriert wichtigen Informationsanteil.
- Urteilsvermögen und Intuition sind hilfreich, aber Bewertung basiert letztendlich auf tatsächlicher Leistung beim Lösen von Probleminstanzen.

Beispiel Routenplanung: Die Straßenentfernung wird durch die Luftlinienentfernung abgeschätzt.



Straight-line distance
to Bucharest

Arad	366
Bucharest	0
Craiova	160
Dobreta	242
Eforie	161
Fagaras	178
Giurgiu	77
Hirsova	151
Iasi	226
Lugoj	244
Mehadia	241
Neamt	234
Oradea	380
Pitesti	98
Rimnicu Vilcea	193
Sibiu	253
Timisoara	329
Urziceni	80
Vaslui	199
Zerind	374

Bei der Bestensuche erfolgt die Expansion eines Knotens auf Basis der heuristischen Funktion:

- Die Heuristik wählt den vielversprechendsten Zustand aus.
- Hierzu werden in der Agenda die Knoten zusammen mit ihrer Bewertung abgelegt.
- Es wird jeweils der Knoten der Agenda expandiert, der die geringste Bewertung aufweist.
- Die Agenda hat also die Form einer Prioritätswarteschlange (priority queue).
- Ansonsten ist die Bestensuche analog zur Tiefen- und Breitensuche.

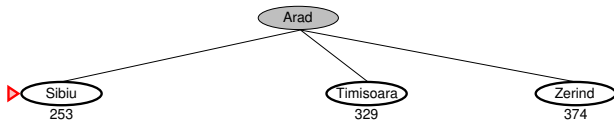
Algorithmus:

```
agenda = new PriorityQueue();
agenda->insert(startknoten, h(startknoten));
closed = new HashMap();
while (!agenda->isEmpty()) {
    akt = agenda->extractMin();
    if (zielknoten(akt))
        return akt;
    if (!closed->contains(akt)) {
        closed->insert(akt);
        foreach node in nachfolger(akt)
            agenda->insert(node, h(node));
    }
}
return 0;
```

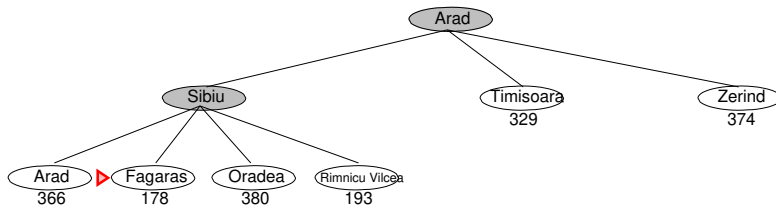
Routenplanung:



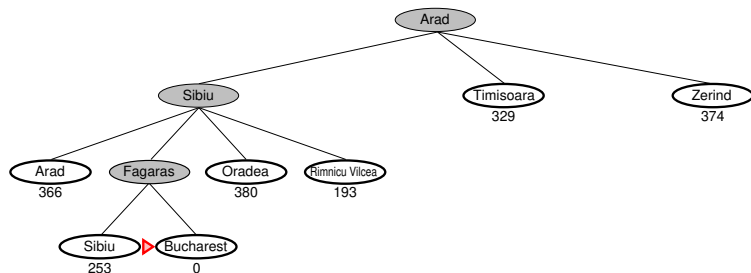
Routenplanung:



Routenplanung:



Rutenplanung:



- Die gefundene Lösung ist nicht optimal!
- Bestensuche wählt Route über Sibiu und Fagaras, obwohl die Route über Rimnicu Vilcea und Pitesti kürzer ist.
- Heuristiken können versagen: Eine Heuristik ist nur eine informierte Vermutung, basierend auf Erfahrung und Intuition.

Die Bestensuche ist anfällig für falsche Startknoten!

- Beispiel: Suche Route von Iasi nach Fagaras.
 - Zuerst wird Neamt expandiert, weil es näher am Ziel ist als Vaslui, aber die Route führt in eine Sackgasse.
 - Die richtige Route führt zunächst über Vaslui, obwohl es weiter vom Ziel entfernt ist als Neamt.
- Heuristik expandiert unnötige Knoten.
- Ohne Erkennung wiederholter Zustände würde die Suche zwischen Iasi und Neamt oszillieren.

Analogie zur Tiefensuche:

- Verfolge nur einen Pfad zum Ziel, um die Speicherkomplexität gering zu halten
- und springe zurück, falls in eine Sackgasse verzweigt wurde.

Gleiche Nachteile wie Tiefensuche:

- Die Bestensuche ist nicht optimal und
- nicht vollständig (kann in unendlichem Pfad hängen bleiben).

Komplexität: Sei b die Verzweigungsrate und m die maximale Tiefe des Baums. Dann ergeben sich folgende Komplexitäten:

- Zeit: $O(b^m)$
- Platz: $O(b^m)$

Eine heuristische Funktion h heißt *fair* gdw. es zu jedem $n \geq 0$ nur endlich viele Knoten s gibt mit $h(s) \leq n$.

- Fairness entspricht der Vollständigkeit bei uninformierten Suchverfahren.
- Ist eine heuristische Funktion fair, so wird ein Zielknoten gefunden, falls ein solcher existiert.

Problem:

- Die Bestensuche vernachlässigt die Kosten bei der Anwendung der Operatoren.
- Wird die Güte einer Lösung charakterisiert durch diese Operatorkosten, so findet die Bestensuche i.Allg. keine optimale Lösung.
 - Routenberechnung: Distanz vom Start zum aktuellen Knoten
 - 8-Puzzle: Anzahl bisher durchgeführter Züge

Heuristiken sind fehlbar, daher verfolgt ein Suchalgorithmus evtl. einen Weg, der nicht zu einem Ziel führt.

- Lösung des Problems bei Tiefensuche: Ein Tiefenzähler wurde genutzt, um aussichtslose Wege zu entdecken.
- Diese Idee ist auch bei heuristischer Suche anwendbar.

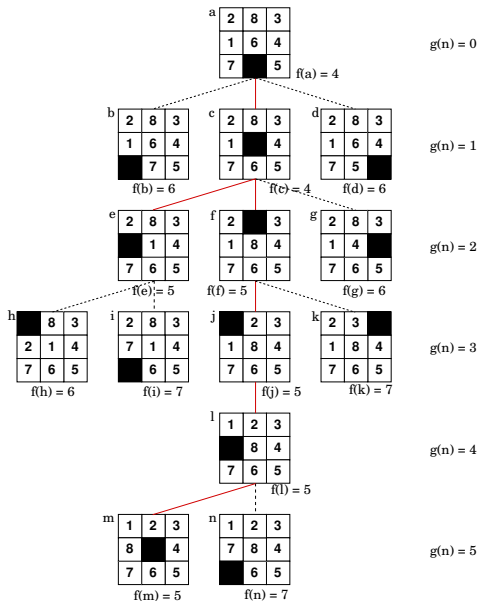
Bei gleicher heuristischer Bewertung ziehe den Zustand vor, der sich näher an der Wurzel des Graphen befindet.

- Dieser Knoten hat eine größere Wahrscheinlichkeit, auf dem *kürzesten Weg* zum Ziel zu liegen.
- Jedem Zustand wird ein Tiefenzähler zugewiesen, der für jede Ebene um eins hochgezählt wird.
- Bewertungsfunktion f ist die Summe aus zwei Komponenten.
 - $g(n)$: Tatsächliche Länge des Weges vom Start zum Zustand n .
 - $h(n)$: Schätzung der Entfernung von n zu einem Ziel.

Algorithmus:

```
agenda = new PriorityQueue();
g(startknoten) := 0;
agenda->insert(startknoten, h(startknoten));
closed = new HashMap();
while (!agenda->isEmpty()) {
    akt = agenda->extractMin();
    if (zielknoten(akt))
        return akt;
    if (!closed->contains(akt)) {
        closed->insert(akt);
        foreach n in nachfolger(akt) {
            g(n) := g(akt) + k(akt, n);
            agenda->insert(n, g(n) + h(n));
        }
    }
}
return 0;
```

Schiebepuzzle



- Die Komponente $g(n)$ verleiht der Suche eine stärkere Tendenz zur Breitensuche.
- Dies verhindert, dass die Suche von einer fehlerhaften Bewertung in die Irre geführt wird.
- Der Algorithmus verliert sich nicht dauerhaft in einem unendlichen Zweig, falls eine Heuristik kontinuierlich gute Bewertungen abgibt.
- Jeder Zustand muss darauf untersucht werden, ob er schon mal aufgetreten ist, um Schleifen zu verhindern.
- Wird der Zustand auf einem kürzeren Weg erreicht, muss die Bewertung angepasst werden.
 - Die Methode `insert` muss prüfen, ob der Zustand n bereits in der Agenda enthalten ist und ggf. den Wert anpassen.
 - Wurde Zustand n bereits abgearbeitet und in die closed list eingefügt, *dann muss n der closed list entnommen und wieder in die Agenda aufgenommen werden.*

Sei $p = (s_0, s_1, \dots, s_r)$ eine Folge von Zuständen und s_{i+1} sei durch Anwendung eines Zustandsübergangsoperators auf s_i erreichbar.

Wenn beim Übergang von s_i nach s_{i+1} Kosten in Höhe von $k(s_i, s_{i+1})$ anfallen, dann seien die Kosten $k(p)$ der Zustandsfolge definiert durch:

$$k(p) := \sum_{i=0}^{r-1} k(s_i, s_{i+1})$$

Für einen Zustand s sei:

- $g^*(s)$ minimale Kosten vom Start nach s , also

$$g^*(s) := \inf\{k(p) \mid p \text{ ist Weg vom Startzustand nach } s\}$$

- $h^*(s)$ minimale Kosten von s zum Ziel, also

$$h^*(s) := \inf\{k(p) \mid p \text{ ist Weg von } s \text{ zu einem Zielzustand}\}$$

- also: Start $\xrightarrow{g^*(s)}$ s $\xrightarrow{h^*(s)}$ Ziel

Problem: Finde eine Zustandsfolge p' vom Startzustand s_0 in einen Zielzustand z , die minimale Kosten aufweist, d.h.:

$$k(p') = h^*(s_0) \text{ bzw.}$$

$$k(p') = \inf\{g^*(z) \mid z \text{ ist Zielzustand}\}$$

Eine heuristische Funktion h heißt *zulässiger Schätzer* bzw. *zulässig*, falls $h(s) \leq h^*(s)$ für alle Zustände s des Zustandsraums gilt. Mit anderen Worten: Ein zulässiger Schätzer darf die echten Kosten nicht überschätzen!

Zulässige Schätzer sind:

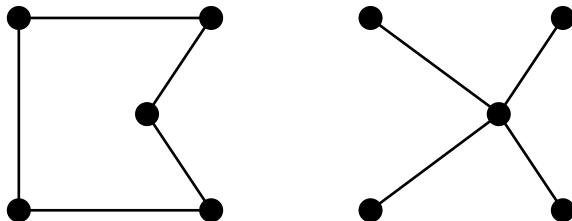
- Die heuristischen Funktionen für das Schiebepuzzle.
- Die Luftlinienentfernung beim Routenproblem.

Bei kombinatorischen Optimierungsproblemen werden als zulässige Schätzer häufig effizient lösbare Relaxationen des Problems verwendet: Vernachlässige Nebenbedingungen, so dass sich ein einfacher zu lösendes Problem ergibt.

Beispiel: Minimaler Spannbaum als Relaxation für die Berechnung eines minimalen Hamiltonschen Weges.

Well-known example: Travelling salesperson problem (TSP).

Find the shortest tour visiting all cities exactly once.



Minimum spanning tree can be computed in $O(n^2)$ and is a lower bound on the shortest (open) tour.

Relaxationen entstehen oft durch Weglassen oder Lockern von Bedingungen.

Beispiel: 8-Puzzle

- Ein Stein kann direkt überall platziert werden, auch auf bereits belegten Feldern.
- h_1 ergibt exakte Anzahl Schritte der optimalen Lösung.
- Ein Stein kann in alle Richtungen bewegt werden, auch auf ein schon belegtes Feld.
- h_2 ergibt exakte Anzahl Schritte der optimalen Lösung.

Achtung: Die Kosten einer optimalen Lösung eines relaxierten Problems sind *zulässige Schätzer* des originalen Problems!

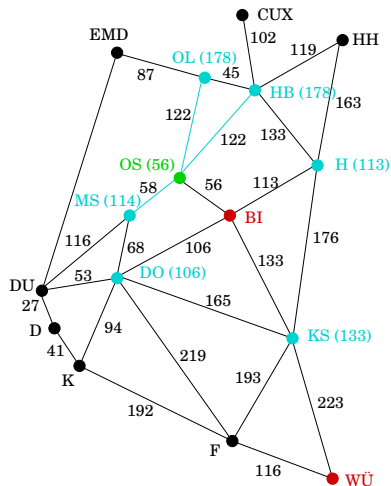
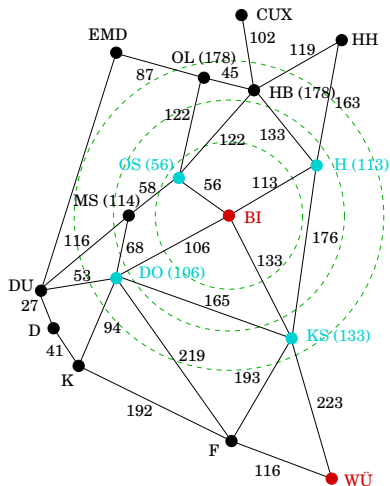
Eine optimale Lösung des Originalproblems ist auch eine Lösung des relaxierten Problems, und ist mindestens so teuer wie die optimale Lösung des relaxierten Problems.

- Durch eine Verringerung von $g(s)$ für einen Zustand s kann auch eine Verringerung von $f(s)$ auftreten.
- Dies kann auch für schon expandierte Knoten der Fall sein.
- Sind Bewertungen von bereits expandierten Knoten anzupassen, *werden diese aus der closed list entfernt und wieder in die Agenda aufgenommen.*
- Dadurch können sie später erneut betrachtet werden.

Monotone Schätzer:

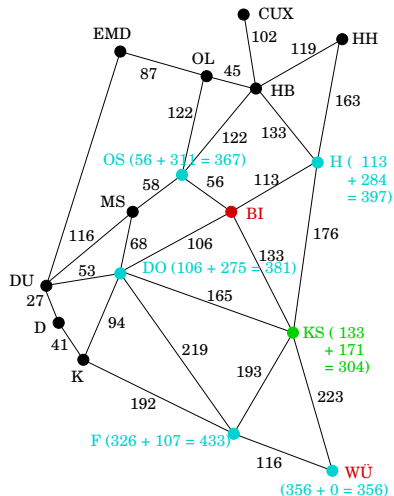
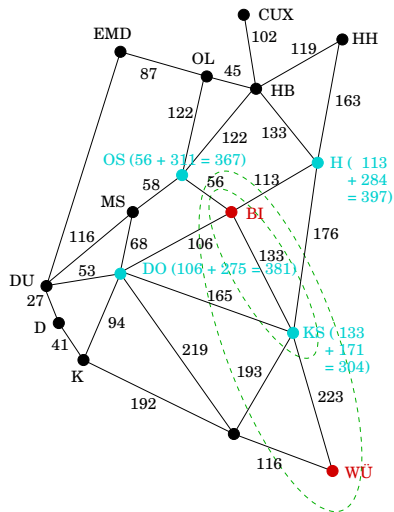
- Zulässiger Schätzer, der die Dreiecksungleichung erfüllt, d.h. für alle Zustände s und alle Nachfolger s' von s gilt:
$$h(s) \leq k(s, s') + h(s')$$
- *Bei Einsatz von monotonen Schätzern werden expandierte Knoten nie wieder betrachtet.*

Dijkstra: Suche von Bielefeld nach Würzburg



Die Suche breitet sich in alle Richtungen gleichmäßig aus, also kreisförmig um den Startknoten, bis das Ziel gefunden wird.

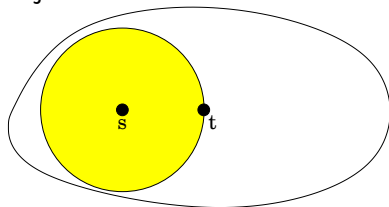
A*: Suche Weg von Bielefeld nach Würzburg



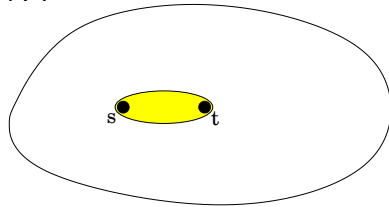
Durch die Schätzfunktion findet die Suche „zielgerichtet“ statt, nicht in zentrischen Kreisen, sondern in Konturen.

Durch die Schätzfunktion findet die Suche „zielgerichtet“ statt, nicht in zentrischen Kreisen, sondern in Konturen.

Dijkstra:



A^* :



Beschleunigung: umso größer, je weiter s und t auseinander liegen

Aus der Vorlesung „Effiziente Algorithmen“ wissen wir, dass der Algorithmus von Dijkstra korrekt ist, solange keine Kante e mit negativer Bewertung $cost(e)$ im Graphen enthalten ist.

```
 $d[s] := 0$   
 $Q.clear()$   
 $Q.add(s, 0)$   
while not  $Q.isEmpty()$  do  
     $u := Q.extractMin()$   
    for each  $e = (u, v) \in E$  do  
        if not  $Q.contains(v)$  or  $d[v] > d[u] + cost(e)$   
        then  $d[v] := d[u] + cost(e)$   
            if  $Q.contains(v)$   
            then  $Q.decreaseKey(v, d[v])$   
            else  $Q.insert(v, d[v])$ 
```

Korrektheit A^* -Algorithmus

Sei h ein monotoner Schätzer, so dass $h(u) \leq \text{cost}(u, v) + h(v)$ für alle $(u, v) \in E$ gilt.

```
d[s] := 0
Q.clear()
Q.add(s, h(s))
while not Q.isEmpty() do
    u := Q.extractMin()
    for each e = (u, v) ∈ E do
        if not Q.contains(v) or d[v] > d[u] + cost(e)
        then d[v] := d[u] + cost(e)
            if Q.contains(v)
            then Q.decreaseKey(v, d[v] + h(v))
            else Q.insert(v, d[v] + h(v))
```

Abarbeitung der Knoten:

- Dijkstra: Knoten u vor Knoten v , wenn $d[u] < d[v]$
- A^* : Knoten u vor Knoten v , wenn $d[u] + h(u) < d[v] + h(v)$

Wir definieren $cost_h(u, v) := cost(u, v) - h(u) + h(v)$ und ersetzen in G die alte durch die neue Kostenfunktion und erhalten so G_h .

- Die Länge aller Pfade vom Start s zum Ziel t ändert sich um den selben Wert $h(t) - h(s)$, daher bleiben alle kürzesten Wege erhalten: $d_h(s, u) = d(s, u) - h(s) + h(u)$
- Die Reihenfolge der besuchten Knoten von Dijkstra im neuen Graphen G_h entspricht also der von A^* in G , da $h(s)$ für alle Knoten gleich ist.
- A^* ist also korrekt, wenn $cost_h(u, v) \geq 0$ für alle $(u, v) \in E$ ist. Dies folgt unmittelbar aus der Dreiecksungleichung $h(u) \leq cost(u, v) + h(v)$ des monotonen Schätzers h .

Übung: Betrachten Sie das Rucksackproblem:

- Gegeben: Eine Menge $G = \{g_1, \dots, g_n\}$ von Gegenständen und ein maximales Gesamtgewicht C .
- Jeder Gegenstand $g \in G$ hat ein Gewicht $w(g)$ und liefert einen Profit $p(g)$.

Bestimme eine Teilmenge $F \subseteq G$, so dass folgendes gilt:

- Der Gesamtprofit $\sum_{g \in F} p(g)$ von F ist maximal
- unter der Nebenbedingung $\sum_{g \in F} w(g) \leq C$.

Zentrale Fragen:

- Formulierung des Rucksackproblems als Suchproblem?
- Was muss ein A*-Algorithmus für das Rucksackproblem berücksichtigen?
- Abschätzung des noch erzielbaren Nutzens?

Übung: Implementieren Sie den A*-Algorithmus zum Lösen des $(n^2 - 1)$ -Puzzle und vergleichen Sie die Laufzeiten für $n = 3, 4, 5$.

Wie effektiv ist die heuristische Funktion h_2 im Vergleich zu h_1 ? Ist die heuristische Funktion h_2 auch für größere Puzzle effektiv?

Vergleichen Sie die Laufzeiten des obigen Programms mit den Laufzeiten des Integer-Programms, das wir mittels GLPK implementiert hatten.

- Motivation
- Uninformierte Suche
- Informierte Suche
- *Lokale Suche*
- Spiele

Bei vielen Problemen ist nicht der Lösungsweg interessant, stattdessen ist der Zielzustand die eigentliche Lösung.

- 8-Damen-Problem: Nur die Positionen der Damen sind wichtig, nicht in welcher Reihenfolge die Damen platziert wurden.
- Entwurf Integrierter Schaltungen: Nur die letztendliche Platzierung der Bauteile und die Verdrahtung ist wichtig, nicht der Weg, wie das Ziel erreicht wurde.
- Problem des Handlungsreisenden: Nur die Tour ist wichtig.
- Stunden-/Fahrplan Erstellung: Nur der Plan ist wichtig.

Trotz Einsatz von Heuristiken ist die Laufzeit der bisher betrachteten Suchalgorithmen oft noch zu schlecht für einen praktischen Einsatz.

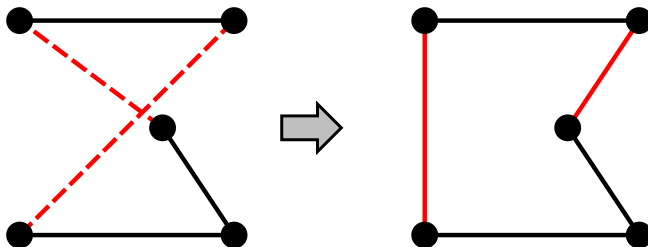
Lokale Suchalgorithmen

- arbeiten mit nur einem einzigen aktuellen Zustand anstelle von mehreren Pfaden.
- versuchen iterativ den aktuellen Zustand zu verbessern.
- besuchen in der Regel nur benachbarte Zustände. Im Gegensatz zu Breiten-, Tiefen- oder Bestensuche, wo die Suche mit dem nächsten Zustand der Agenda fortgesetzt wird.
- speichern die Pfade zur Problemlösung nicht ab.

Vorteile:

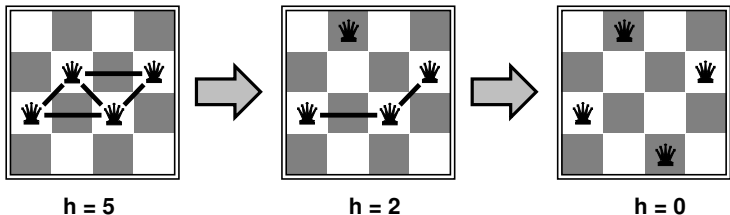
- Sehr geringer Speicherverbrauch.
- Oft werden vernünftige Lösungen in riesigen Suchräumen gefunden.

traveling salesperson problem TSP: start with any complete tour, perform pairwise exchanges



variants of this approach get within 1% of optimal very quickly with thousands of cities

n -queens problem: move a queen to reduce number of conflicts



almost always solves n -queens problems almost instantaneously for very large n , e.g., $n = 1$ million

number of pairs of queens that are attacking each other, either directly or indirectly

18	12	14	13	13	12	14	14
14	16	13	15	12	14	12	16
14	12	18	13	15	12	14	14
15	14	14		13	16	13	16
	14	17	15		14	16	16
17		16	18	15		15	
18	14		15	15	14		16
14	14	13	17	12	14	12	18

choose randomly among the set of best successors, if there is more than one

Hill Climbing Search

sometimes called *greedy local search*, because it grabs a good neighbor state without thinking ahead about where to go next

```
function HILL-CLIMBING(problem) returns a state that is a local maximum
  inputs: problem, a problem
  local variables: current, a node
                  neighbor, a node

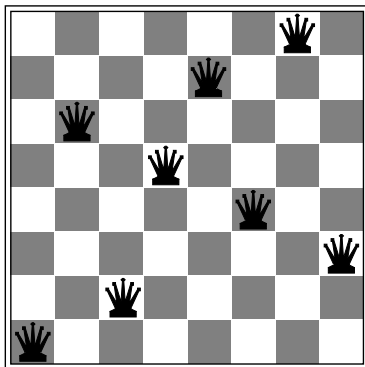
  current ← MAKE-NODE(INITIAL-STATE[problem])
  neighbor ← a highest-valued successor of current
  while VALUE[neighbor] ≤ VALUE[current] do
    current ← neighbor
    neighbor ← a highest-valued successor of current
  end
  return STATE[current]
```

steepest ascent version

“it’s like climbing Everest in thick fog with amnesia”

Hill Climbing Search

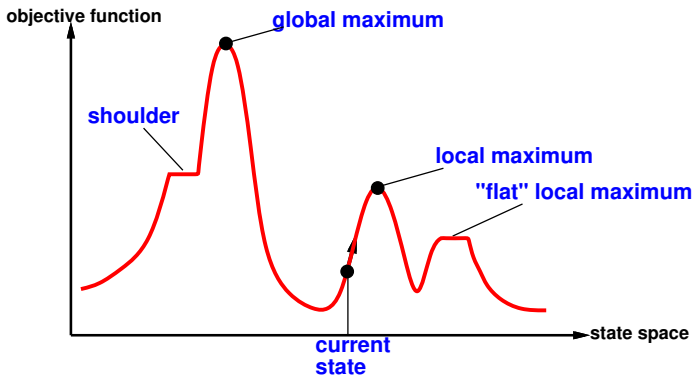
search may be lost in local minimum



each move of a queen would increase the number of conflicts

Hill Climbing Search

useful to consider state space landscape



NP-complete problems typically have an exponential number of local maxima to get stuck on

Variants of Hill Climbing

- *random sideways moves*
 - ☺ escape from shoulders → comes at a cost of higher running time
 - ☹ loop on flat maxima → put a limit on the number of consecutive sideways moves allowed
- *stochastic* choose at random from among the uphill moves
 - ☺ in some landscapes it finds better solutions than steepest ascent
 - ☹ usually converges more slowly
- *first-choice* implements stochastic hill climbing by generating successors randomly until one is generated that is better than the current state
 - good strategy when a state has many of successors
- *random-restart* a reasonably good local maximum can often be found after a small number of restarts
 - overcomes local maxima — trivially complete

Idea: keep k states instead of 1; choose top k of all their successors

- Not the same as k searches run in parallel! Searches that find good states recruit other searches to join them.
- useful information is passed among the k parallel search threads
- Problem: quite often, all k states end up on same local hill

Idea: choose k successors randomly, biased towards good ones

→ stochastic beam search

- observe the close analogy to natural selection

Idee der Tabu-Suche:

- Um lokale Optima verlassen zu können, wird in jedem Schritt die beste erlaubte Nachbarlösung gewählt, auch wenn diese schlechter ist als die aktuelle. → Best-Improvement-Methode
- Nutze eine Gedächtnisstruktur, um die Suche zu steuern.

Um Zyklen beim Traversieren des Suchraums zu vermeiden, werden bereits besuchte Knoten in einer Liste gespeichert.

- Die in der Liste enthaltenen Knoten dürfen in den nächsten Schritten nicht ausgewählt werden.
- Verbiete den Knoten so lange, bis das Wiedererreichen dieser Lösung unwahrscheinlich ist. → Kurzzeitgedächtnis
- Probleme, die experimentell gelöst werden müssen:
 - Wie groß soll die Liste sein?
 - Nach wievielen Schritten soll ein Eintrag entfernt werden?

Großer Speicherbedarf, wenn ganze Knoten in der Tabu-Liste gespeichert werden. Daher werden oft nur Attribute verboten:

- Traveling Salesperson: verbiete den Tausch zweier Städte
- Färbung von Graphen: verbiete Farbzuzuweisung i an Knoten k

Ein Langzeitgedächtnis kann das Verfahren weiter verbessern:

- Mit der Information, wie oft Lösungen untersucht wurden, kann die Suche so gesteuert werden, das bisher wenig betrachtete Bereiche des Lösungsraums untersucht werden.
- Problem: Speicherbedarf wird noch größer.

Simulated Annealing Search

Idea: escape local maxima by allowing some “bad” moves but gradually decrease their size and frequency

- devised by Metropolis et al., 1953, for physical process modelling
- widely used in VLSI layout, airline scheduling, etc.

Simulated Annealing Search

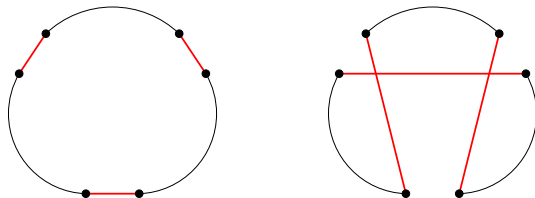
```
function SIMULATED-ANNEALING(problem, schedule) returns a solution state
  inputs: problem, a problem
           schedule, a mapping from time to "temperature"
  local variables: current, a node
                   next, a node
                   T, a "temperature" controlling prob. of downward steps

  current  $\leftarrow$  MAKE-NODE(INITIAL-STATE[problem])
  for t  $\leftarrow$  1 to  $\infty$  do
    T  $\leftarrow$  schedule[t]
    if T = 0 then
      return current
    next  $\leftarrow$  a randomly selected successor of current
     $\Delta E \leftarrow$  VALUE[next] - VALUE[current]
    if  $\Delta E > 0$  then
      current  $\leftarrow$  next
    else current  $\leftarrow$  next only with probability  $e^{\Delta E / T}$ 
  end
```

Exercise: Solve instances of n-queens and TSP by hill climbing, tabu search, local beam search and simulated annealing.

Measure the search cost and the percentage of solved problems and graph these against the optimal solution cost. Comment on your results.

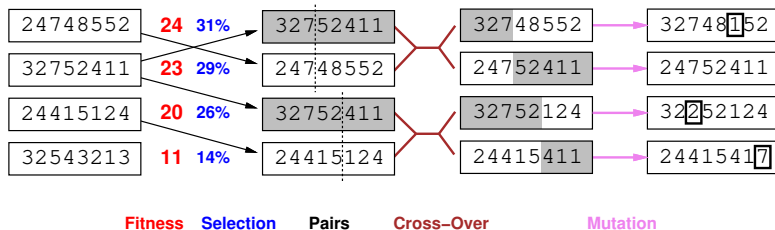
TSPLIB is a library of sample instances for the TSP. Instances of symmetric traveling salesman problem are available. Use 2-opt and 3-opt on these instances and compare the results.



stochastic beam search + generate successors from pairs of states

- GAs begin with a set of k randomly generated states, called the *population*
- each state, or *individual*, is represented as a string over a finite alphabet
- each state is rated by the evaluation function or *fitness function*
- a fitness function should return higher values for better states
for the 8-queens problem we use the number of nonattacking pairs of queens, which has a value of 28 for a solution
- in the following example the probability of being chosen for reproducing is directly proportional to the fitness score

8-queens problem



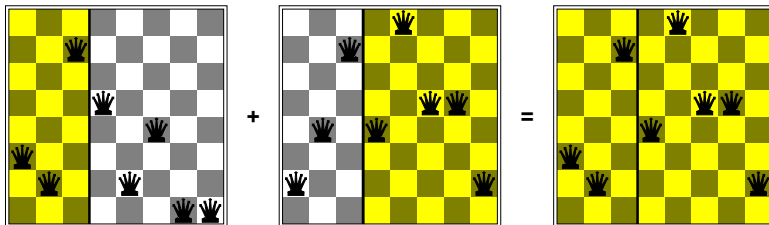
- a random choice of two pairs is selected for reproduction, in accordance with the shown probabilities
- a variant of this selection rule, called *culling*, in which all individuals below a given threshold are discarded, can be shown to converge faster than the random version
- for each pair to be mated, a *crossover* point is randomly chosen from the positions in the string

- finally, each location is subject to random mutation with a small independent probability
in the 8-queens problem, this corresponds to choosing a queen at random and moving it to a random square in its column
- primary advantage, if any, of genetic algorithms comes from the crossover operation: combine large blocks of letters that have evolved independently to perform useful functions, thus raising the level of granularity at which the search operates

GAs require states encoded as strings (GPs use programs)

- an 8-queens state must specify the positions of 8 queens, each in a column of 8 squares, and so requires $8 \times \log_2(8) = 24$ bits
- alternatively, the state could be represented as 8 digits, each in the range from 1 to 8
- behave the two encodings differently?

Crossover helps iff substrings are meaningful components



- *schema*: a substring in which some of the positions can be left unspecified, for example, 246*****
- strings that match the schema such as 24613578 are called *instances* of the schema
- if the average fitness of the instances of a schema is above the mean, then the number of instances of the schema within the population will grow over time

- successful use of GAs require careful engineering of the representation
- at present, it is not clear whether the appeal of GAs arises from their performance or from their pleasing origins in the theory of evolution

- Motivation
- Uninformierte Suche
- Informierte Suche
- Lokale Suche
- *Spiele*

Was fasziniert uns an Spielen?

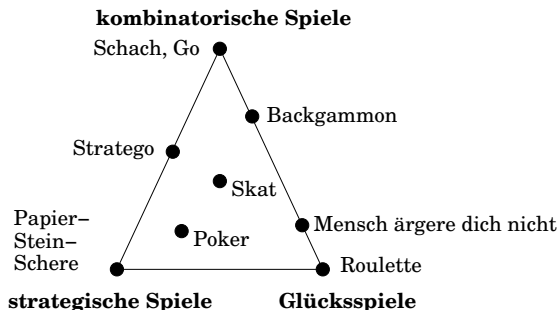
- Vor allem der ungewisse Ausgang des Spiels.

Diese Ungewissheit basiert auf

- *Zufall*: Tritt z.B. durch Würfeln (Backgammon) oder Mischen von Spielkarten (Skat) auf.
- *Kombinatorik*: Spielregeln legen die Handlungsmöglichkeiten der Spieler fest, z.B. bei Schach, Go oder Vier gewinnt.
- *mangelnder Information*: Da man z.B. nicht die Karten der Mitspieler oder deren Absichten kennt, erhält man nur ein Teilkenntnis des erreichten Spielstandes.

Arten von Spielen

- *Glücksspiel*: Der Einfluss des Zufalls dominiert gegenüber denen der Spieler.
- *Kombinatorisches Spiel*: Ungewissheit beruht auf den vielfältigen, praktisch unüberschaubaren Zugmöglichkeiten.
- *Strategisches Spiel*: Ungewisser Ausgang durch imperfekte Information.



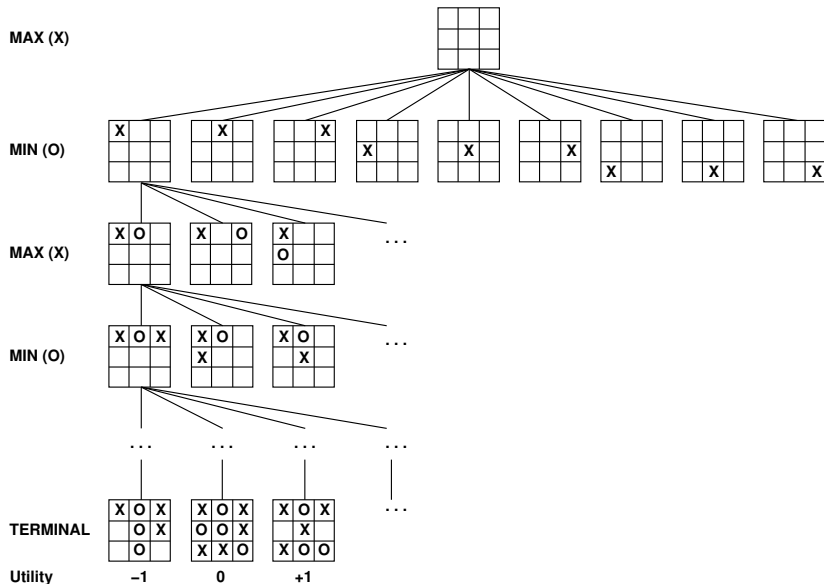
- *Zug*: Spielabschnitt, der genau eine durch die Spielregeln festgelegte Handlungsmöglichkeit eines Spielers umfasst. Wird manchmal auch Halbzug genannt.
- *Zugzwang-Phänomen*: Wenn ein Spieler nicht passen kann, verschlechtert er evtl. seine Situation durch jeden ihm möglichen Zug.
- *Kooperative Züge* stärken beide Spieler gleichzeitig (selten).
- *Nullsummenspiel*: Normalerweise fördert ein Zug den einen Spieler genauso stark, wie er den anderen behindert – die Summe der Nutzeneffekte beider Spieler ist immer gleich Null.

Kombinatorische Spiele als Suchprobleme

- Es gibt genau zwei Spieler: Max und Min.
- Max hat immer den ersten Zug.
- Es ist eine Anfangsstellung des Spiels definiert.
- Jeder Spieler besitzt die vollständige Information über die Zugmöglichkeiten des Gegners.
- Die Spieler ziehen abwechselnd, es ist kein Aussetzen möglich.
- Nach endlich vielen Zügen wird stets eine Spielstellung ohne Fortsetzungsmöglichkeiten erreicht.
- Gesucht ist eine Gewinn-Strategie für Max.
- Eine Nutzenfunktion bewertet die Endzustände und damit den Ausgang des Spiels numerisch.

win	+1	Max hat gewonnen
draw	0	unentschieden
loss	-1	Min hat gewonnen

Suchbaum für Tic-Tac-Toe



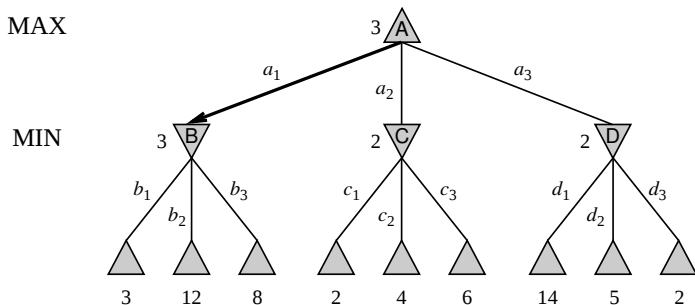
Rational agierende Spieler

- denken Zug für Zug voraus, abwechselnd einen eigenen und einen des Gegners.
- weisen den erreichbaren Positionen Werte zu: Gewinn für mich, unentschieden oder Gewinn für den Gegner.
- versuchen in jedem Zug ihren Vorteil zu maximieren, während der Gegner diesen minimieren will.

Minimax-Algorithmus

Welchen Zug muss Max machen, um zu gewinnen?

- Generiere den vollständigen Suchbaum für das Spiel.
- Wende die Nutzenfunktion auf jeden Endzustand an.
- Weise den Knoten im Suchbaum wie folgt Werte zu:
 - Ein Knoten von Max erhält als Wert das Maximum der Kinder.
 - Ein Knoten von Min erhält als Wert das Minimum der Kinder.
- Max wählt einen Zug passend zum Wert des Wurzelknotens



Minimax-Algorithmus

function MINIMAX-DECISION(*state*) *returns an action*

inputs: *state*, current state in game

$v \leftarrow \text{MAX-VALUE}(\textit{state})$

return the *action* in **SUCCESSORS**(*state*) with value v

function MAX-VALUE(*state*) *returns a utility value*

if **TERMINAL-TEST**(*state*)

then return **UTILITY**(*state*)

$v \leftarrow -\infty$

foreach s in **SUCCESSORS**(*state*) **do**

$v \leftarrow \text{MAX}(v, \text{MIN-VALUE}(s))$

return v

function MIN-VALUE(*state*) *returns a utility value*

if **TERMINAL-TEST**(*state*)

then return **UTILITY**(*state*)

$v \leftarrow \infty$

foreach s in **SUCCESSORS**(*state*) **do**

$v \leftarrow \text{MIN}(v, \text{MAX-VALUE}(s))$

return v

Beispiel: Nimm-Spiel

- Das Spiel beginnt mit einem Haufen von n Spielmarken.
- Die Spieler ziehen abwechselnd.
- Bei jedem Zug muss ein Spieler ein Häufchen Spielmarken in nicht-leere unterschiedlich große Häufchen teilen.
- Der erste Spieler, der nicht mehr ziehen kann, verliert.

Spielbaum? Minimax-Verfahren?

Eine erschöpfende Suche bis zu den Endknoten ist in der Regel nicht möglich:

- Time complexity: $O(b^m)$
 - Space complexity: $O(b \cdot m)$ (depth-first exploration)
 - for chess, $b \approx 35$, $m \approx 100$ for “reasonable” games
- exact solution completely infeasible

- Durchsuche den Zustandsraum nur bis zu einer vordefinierten Anzahl n von Ebenen.
- Um den Wert von n festzulegen, ist der Ressourcenverbrauch für Zeit und Platz zu berücksichtigen.
- Die zu bewertenden Zustände sind dann aber in der Regel keine Endzustände.

Daher ist eine heuristische Bewertungsfunktion anzuwenden.

- Der Wert am Wurzelknoten zeigt nicht mehr an, ob das Spiel gewonnen wird!

Es handelt sich nur um den Wert des am höchsten bewerteten Zustandes, der in n Zügen vom Startknoten aus mit Sicherheit erreichbar ist.

- Diese Strategie heißt Vorausschau über n Züge.

Heuristische Bewertungsfunktion

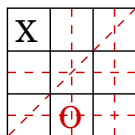
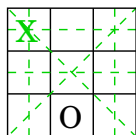
Eine heuristische Bewertungsfunktion enthält Wissen über das Spiel, um eine Stellung zu bewerten.

Beispiel: Heuristische Bewertungsfunktion für Tic-Tac-Toe.

$$Eval(z) = W(z) - L(z)$$

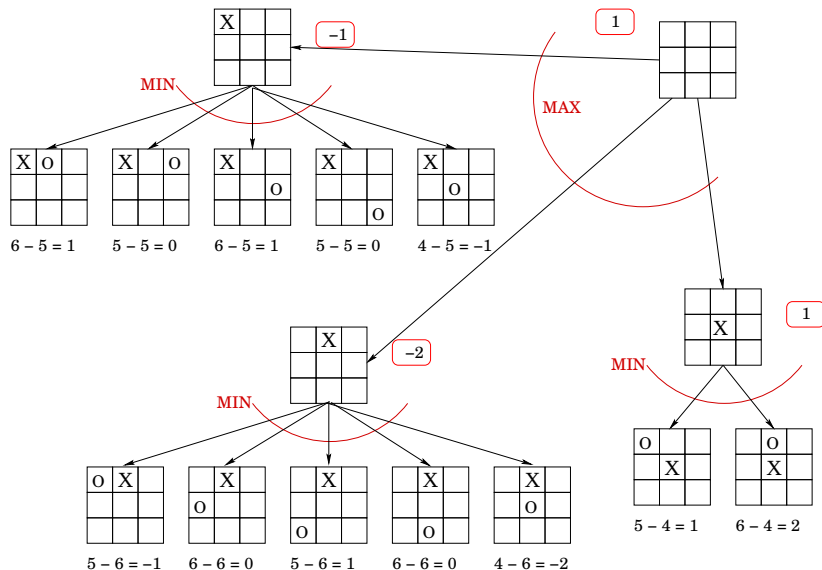
mit

- $W(z)$: Anzahl eigener Gewinnmöglichkeiten (Win).
- $L(z)$: Anzahl gegnerischer Gewinnmöglichkeiten (Loss).
- $Eval(z)$: Bewertung des Zustandes.



- X hat 6 mögliche Gewinnwege
- O hat 5 mögliche Gewinnwege

Heuristische Bewertungsfunktion



Eine Bewertungsfunktion kann auch aus n einzelnen Merkmalen bestehen, die durch eine gewichtete Summe aggregiert werden:

$$Eval(z) = w_1 \cdot f_1(z) + w_2 \cdot f_2(z) + \dots + w_n \cdot f_n(z)$$

Die Koeffizienten $w_1, \dots, w_n$ legen fest, welche Eigenschaften wichtiger sind als andere.

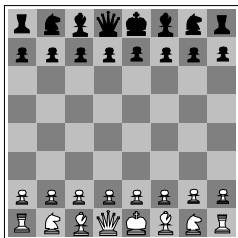
Beispiel Schach:

- $f_1(z)$ Vorteile und Position einer Figur.
- $f_2(s)$ Vorteile und Position einer anderen Figur.
- $f_3(s)$ Kontrolle über das Spielbrett usw.

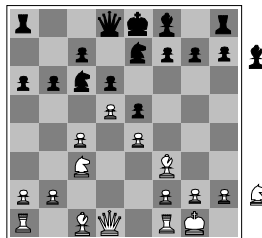
Ist der Vorteil einer Dame größer als bspw. die Kontrolle über das Spielbrett, so wird der entsprechende Koeffizient erhöht.

Führt das Bewertungspolynom zu einer schlechten Zugfolge, müssen die Koeffizienten angepasst werden. → Lernen

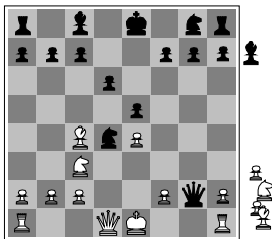
Heuristische Bewertungsfunktion



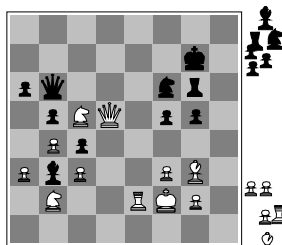
(a) White to move
Fairly even



(b) Black to move
White slightly better



(c) White to move
Black winning



(d) Black to move
White about to lose

Horizont-Effekt

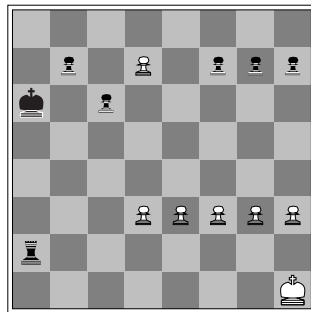
Ein Suchalgorithmus kann in gewissen Situationen zu kurzsichtig sein.

Hier: Materialvorteil für Schwarz, aber Weiß erhält eine Dame!

Besonders interessant, falls der Gegner einen für sich ungünstigen Spielzustand in wenigen Zügen jenseits der Suchtiefe in einen guten Zustand bringen kann.

Bei begrenzter Suchtiefe kann auf der letzten Ebene das Schlagen eines gedeckten Bauern durch eine Dame günstig erscheinen, da es einen Materialvorteil bringt.

Lösung: Shannon definiert Ruhesuche, die noch heute wichtiger Bestandteil jedes Schachprogramms ist.



Black to move

Zur Abschwächung des Horizontproblems wird die Ruhesuche eingesetzt.

- Eine *ruhige Position* ist dann erreicht, wenn der Gegner mit dem nächsten Zug nur eine geringe Änderung des Schätzwertes erzielen kann.

Modifiziere die Minimax-Suche mit beschränkter Tiefe:

- Relevante Blätter im Suchbaum werden durch eine Ruhesuche weitergehend analysiert.
- Zu dem so erweiterten Suchbaum wird mittels Minimax der optimale Zug bestimmt.
- Einen generellen Lösungsansatz zur Eliminierung des Horizonteffektes gibt es bisher nicht.

Ist es wünschenswert, möglichst viele Züge voraus zu schauen?

- Die Bewertungen, die sehr tief im Baum stattfinden, können durch die bloße Tiefe negativ beeinflusst werden.
 - Eine Schätzung von Minimax (das wünschen wir uns) ist etwas anderes als das Minimax von Schätzungen (das tun wir).
- Eine tiefere Suche mit Bewertung und Minimax bedeutet nicht unbedingt, dass dies die bessere Suche ist.

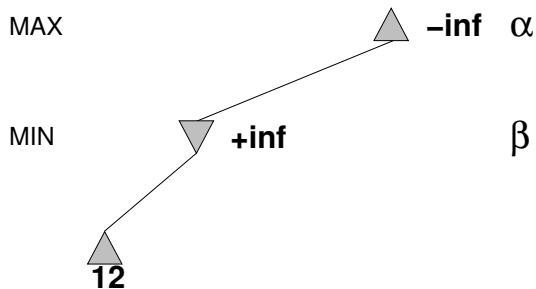
Trotzdem wollen wir versuchen, mittels einer Branch-and-Bound-Technik die Suchtiefe zu steigern. Teilbäume, die kein besseres Ergebnis als das bisher beste liefern können, werden abgeschnitten.

- Für die Berechnung der Minimax-Bewertung irrelevante Zweige werden nicht weiter untersucht.
- Wir benötigen zunächst die Bewertungen der Zustände in der maximalen Suchtiefe, also setzen wir Tiefensuche ein.
- Lege vorläufige Bewertungen fest:
 - Alpha-Werte sind Werte an Max-Knoten.
 - Alpha-Werte können niemals kleiner werden!
 - Beta-Werte sind Werte an Min-Knoten.
 - Beta-Werte können niemals größer werden!
- Diese Werte steuern das vorzeitige Beenden der Suche in einem Teilbaum.
- Wir können die Suche unterhalb eines Max-Knotens s vorzeitig beenden, wenn gilt:

$$\exists \text{Min-Vorgänger } s' \text{ von } s : \alpha(s) \geq \beta(s')$$

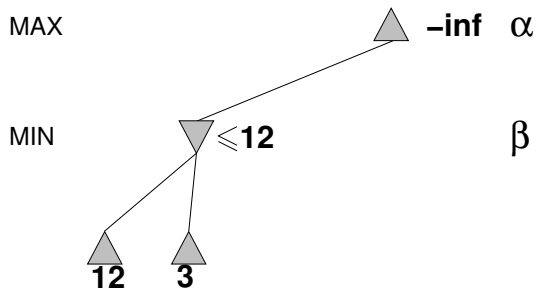
- analoge Definition für Min-Knoten (siehe nächste Folien)

Alpha-Beta-Suche



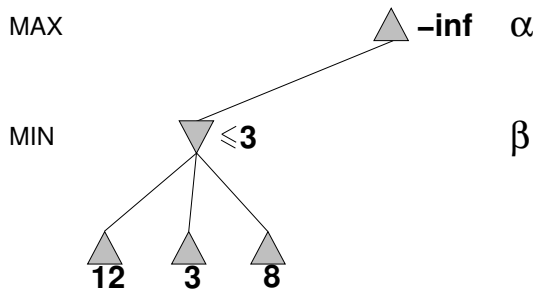
β -Werte können nicht größer werden!

Alpha-Beta-Suche



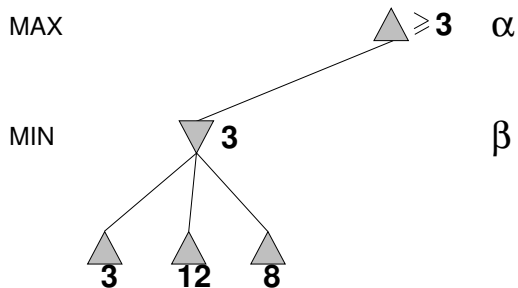
β -Werte können nicht größer werden!

Alpha-Beta-Suche

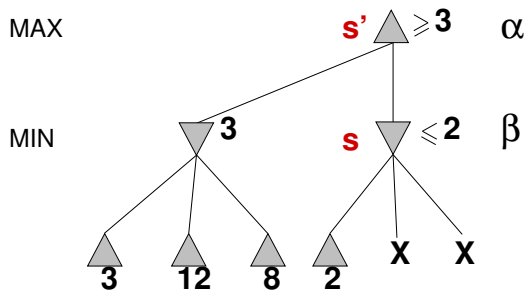


β -Werte können nicht größer werden!

Alpha-Beta-Suche



α -Werte können nicht kleiner werden!

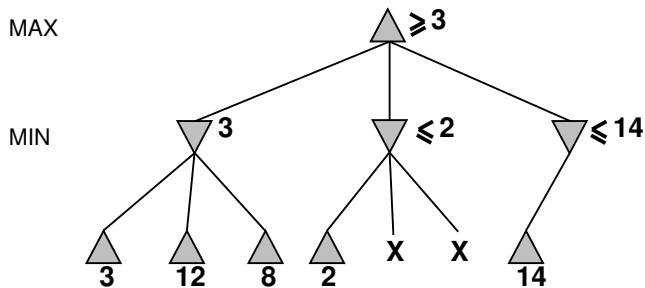


Suche unterhalb des Min-Knotens s vorzeitig beenden, wenn gilt:

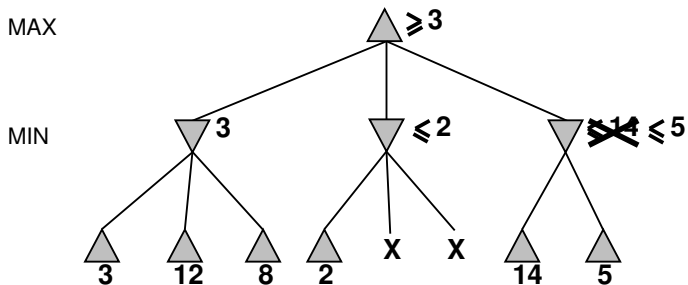
$$\exists \text{Max-Vorgänger } s' \text{ von } s : \beta(s) \leq \alpha(s')$$

Im Zustand s' wird ein Maximum gebildet, das bereits jetzt größer als der Wert $\beta(s)$ ist. Da $\beta(s)$ nicht größer als 2 werden kann, wird die Bearbeitung des Teilbaums unter s abgebrochen.

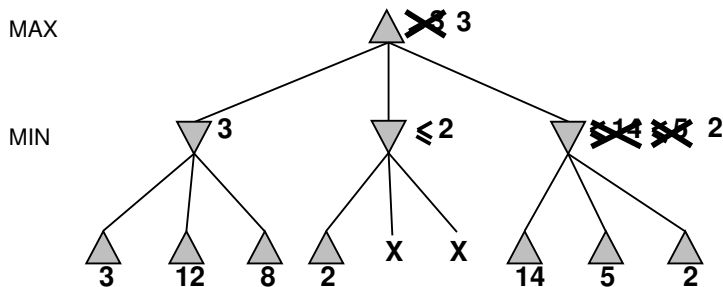
Alpha-Beta-Suche



Alpha-Beta-Suche



Alpha-Beta-Suche



Offensichtlich spielt die Reihenfolge, in der Kindknoten bearbeitet werden, eine wesentliche Rolle. Daher ist es wichtig, zuerst die Züge zu betrachten, die das beste Ergebnis versprechen.

Setze Heuristik ein: Sortiere beim Schach die Züge danach, welche Figur geschlagen wird, oder welche Figur schlägt. „Turm schlägt Dame“ vor „Bauer schlägt Turm“ vor „Turm schlägt Turm“.

function ALPHA-BETA-DECISION(*state*) **returns** an action

inputs: *state*, current state in game

$v \leftarrow \text{MAX-VALUE}(\textit{state}, -\infty, +\infty)$

return the *action* in **SUCCESSORS**(*state*) with value v

function MAX-VALUE(*state*, α , β) **returns** a utility value

inputs: *state*, current state in game

α , the value of the best alternative for MAX along the path to *state*

β , the value of the best alternative for MIN along the path to *state*

if **TERMINAL-TEST**(*state*) **then return** **UTILITY**(*state*)

$v \leftarrow -\infty$

foreach s in **SUCCESSORS**(*state*) **do**

$v \leftarrow \text{MAX}(v, \text{MIN-VALUE}(s, \alpha, \beta))$

if $v \geq \beta$ **then return** v

$\alpha \leftarrow \text{MAX}(\alpha, v)$

return v

function MIN-VALUE(*state*, α , β) **returns** a utility value

same as MAX-VALUE but with roles of α , β reversed

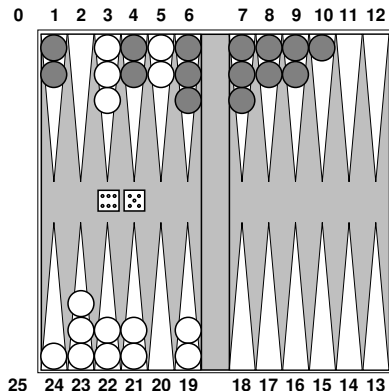
Übung: Implementieren Sie den Minimax-Algorithmus und die α/β -Suche für ein Spiel wie Vier-Gewinnt oder Reversi/Othello. Vergleichen Sie die Laufzeiten beider Verfahren bei gleicher Suchtiefe.

Kombinatorische Spiele mit Zufallselementen

engl.: nondeterministic games

Zufall in Spielen bspw. durch Würfeln (Backgammon) oder durch Mischen und unbekannter Kartenverteilung (Skat).

Im Gegensatz zu rein kombinatorischen Spielen kann eine gute Stellung durch ungünstiges Würfeln oder eine schlechte Kartenverteilung nutzlos sein.



Können wir das Minimax-Prinzip auf solche Spiele übertragen?

Was ändert sich durch das Einbringen von Zufallselementen?

- Ein Spieler kann nicht mehr von einem garantierten Nutzen einer Strategie ausgehen.

Die angewendete Strategie ist oft gut, kann aber im Einzelfall, also z.B. bei ungünstigem Würfeln fehlschlagen.

Wie gehen wir damit um?

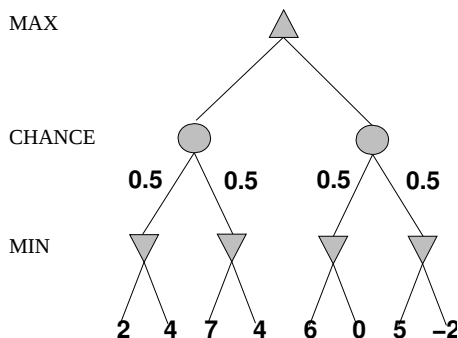
- Bewerte eine Strategie oder auch einen einzelnen Zug mit dem durchschnittlichen Nutzen: Erwartungswert
- Ein Zufallsereignis $\mathcal{X}$ kann zu den Gewinnen $x_1, \dots, x_n$ führen.
- $P(\mathcal{X} = x_i)$: Wahrscheinlichkeit, mit der Gewinn x_i auftritt.
- Erwartungswert $E(\mathcal{X})$ des Zufallsereignisses $\mathcal{X}$:

$$E(\mathcal{X}) = \sum_{i=1}^n x_i \cdot P(\mathcal{X} = x_i)$$

Kombinatorische Spiele mit Zufallselementen

Neben den möglichen Zügen müssen wir nun im Spielbaum für die Bewertung eines Zuges die Zufallskomponente berücksichtigen:

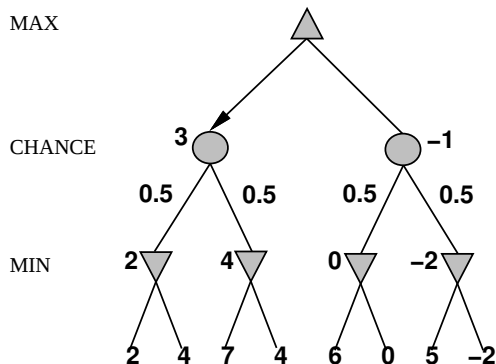
- Zufallsereignisse werden als Knoten repräsentiert.
- Von einem Zufallsknoten ausgehende Kanten sind mit der Wahrscheinlichkeit markiert, mit der der jeweilige Nachfolgezustand erreicht wird.



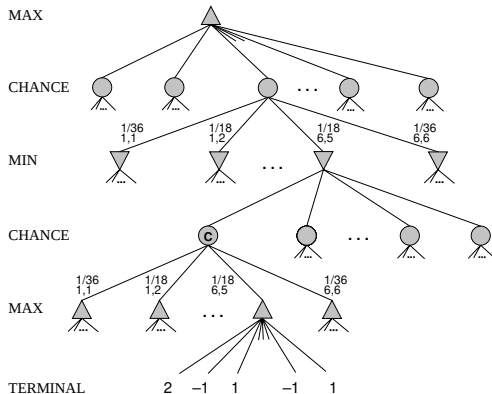
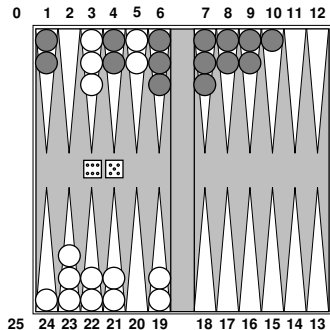
Kombinatorische Spiele mit Zufallselementen

An Zufallsknoten berechnen wir den zugehörigen Erwartungswert:

- Die x_i entsprechen den Bewertungen der Kinder,
- die Wahrscheinlichkeiten $P(\mathcal{X} = x_i)$ ergeben sich durch die Markierungen der Kanten.



Kombinatorische Spiele mit Zufallselementen



- Geschichte und Anwendungen
- Zustandsraumsuche
- *Aussagen- und Prädikatenlogik*
- Einführung in Prolog

Aussagenlogik

- *Grundbegriffe*
- Äquivalenz und Normalformen
- Hornformeln
- Resolution

Unter einer *Aussage* versteht man ein atomares sprachliches Gebilde, das einen der beiden *Wahrheitswerte* wahr oder falsch hat.

Wahre Aussagen:

- Krefeld liegt am Rhein.
- Es gibt unendlich viele Primzahlen.
- $3 + 4 = 7$

Falsche Aussagen:

- Duisburg liegt an der Elbe.
- Die Erde ist eine Scheibe.
- $3 + 4 = 8$

Aussagen, von denen wir zum jetzigen Zeitpunkt nicht wissen, ob sie wahr oder falsch sind:

- Es gibt außerirdisches Leben.
- $P \neq NP$

Es wird letztlich ignoriert, welche inhaltlichen Bedeutungen vorliegen, unser Interesse wird auf die Wahrheitswerte reduziert:

- A : Otto wird krank
- B : der Arzt verschreibt Otto eine Medizin

In der Umgangssprache unterscheiden wir:

A *und* B Otto wird krank *und* der Arzt verschreibt Otto eine Medizin

B *und* A der Arzt verschreibt Otto eine Medizin *und* Otto wird krank

Von solchen Feinheiten befreien wir uns im Folgenden.

Gegenstand der Aussagenlogik:

- Untersuche einfache Verknüpfungen wie *und*, *oder* und *nicht*, zwischen Aussagen und
- lege fest, wie sich Wahrheitswerte der atomaren Bestandteile zu Wahrheitswerten von komplizierteren sprachlichen Gebilden fortsetzen lassen.

Inferenz:

- Herleiten von neuem Wissen auf Basis eines Kalküls.
- Definition des Folgerungsbegriffs.
- Übertragen der semantischen Folgerung auf äquivalente syntaktische Umformungen.

- *Prognosen bzw. logische Ableitungen erstellen:*
Wenn Fakten und Regeln gegeben sind, welche Folgerungen kann man daraus ziehen?
- *Erklärungen finden:*
Wie lässt sich ein Fakt mittels der Regeln erklären?
- *Hypothesen prüfen:*
Können aus den Fakten und Regeln die Hypothesen hergeleitet werden?

Es gibt verschiedene Arten von Schlussfolgerungen²⁰.

Deduktion: Schließe vom allgemeinen Fall auf einen speziellen Fall. Ist immer mit der Vorstellung verbunden, dass es sich um korrekte Schlussfolgerungen handelt, dass das abgeleitete Wissen also stets wahr ist.

Beispiele:

Vögel können fliegen.

Max ist ein Vogel.

Max kann fliegen.

Alle Holländer essen gern Käse.

Hans ist Holländer.

Hans isst gerne Käse.

²⁰Schlussfolgerung wird in der Informatik meist als Inferenz bezeichnet.

Induktion: Man erschließt regelhaftes Wissen aus einzelnen Sachverhalten. Nur so kann man sich weiträumig anwendbare Tatsachen über die wirkliche Welt erschließen. Sonst hätten wir nur Millionen von Einzelerfahrungen.

Allerdings geben wir dabei die Vorstellung auf, das neu abgeleitete Wissen auch stets wahr ist.

Beispiel:

Bello ist ein Hund, Bello bellt, Bello beißt nicht.

Trixi ist ein Hund, Trixi bellt, Trixi beißt nicht.

Waldi ist ein Hund, Waldi bellt, Waldi beißt nicht.

Hunde, die bellen, beißen nicht. (???)

Abduktion: Wir suchen nach einer Erklärung für unsere Beobachtungen, wir bilden Hypothesen. Alle Diagnoseverfahren in der Medizin oder Technik basieren auf Abduktion.

Auch hier ist das abgeleitete Wissen nicht notwendig wahr.

Beispiel:

Max hat lange Ohren.

Hasen haben lange Ohren.

Max ist ein Hase. (???)

Beobachtung

Wissen

Schlussfolgerung

Jeder wissenschaftliche Erkenntnisprozess entsteht aus dem Zusammenspiel von Abduktion, Deduktion und Induktion:

- Finde eine Hypothese (Abduktion).
- Leite Voraussagen aus der Hypothese ab (Deduktion).
- Suche nach Fakten, die die Vorannahmen „verifizieren“ (Induktion).
- Werden keine solche Fakten gefunden, zurück zum Anfang.

- Alle Sätze sind Zeichenketten aus Buchstaben eines Alphabets, die nach präzisen Regeln einer Grammatik angeordnet werden.
- Um eine formale Sprache wie das Resolutionskalkül nutzen zu können, müssen wir sowohl deren Syntax als auch deren Semantik kennen.

Lesen Sie laut den Satz: „Die Aussage $A \wedge B$ ist genau dann wahr, wenn die Aussagen A und B wahr sind.“

- Eine Trennung der Sprachebenen ist unerlässlich.
- Die formale Aussage $A \wedge B$ soll dadurch erklärt werden, dass auf einer metasprachlichen Ebene über die Aussage A wie auch über die Aussage B geredet wird.

Wir bezeichnen mit A_i eine atomare Formel. Eine *Formel* wird durch folgenden induktiven Prozess formuliert:

- Alle atomaren Formeln sind Formeln.
- *Konjunktion*: Für alle Formeln F und G ist $(F \wedge G)$ eine Formel.
- *Disjunktion*: Für alle Formeln F und G ist $(F \vee G)$ eine Formel.
- *Negation*: Für jede Formel F ist $\neg F$ eine Formel.

Eine Formel F , die als Teil einer Formel G auftritt, heißt *Teilformel*.

Beispiel: Die Teilformeln von $F = \neg((A_5 \wedge A_6) \vee \neg A_3)$ sind:

$$F, ((A_5 \wedge A_6) \vee \neg A_3), (A_5 \wedge A_6), A_5, A_6, \neg A_3, A_3$$

Wir führen abkürzende Schreibweisen ein:

$$(F_1 \rightarrow F_2) \quad \text{statt} \quad (\neg F_1 \vee F_2)$$

$$(F_1 \leftrightarrow F_2) \quad \text{statt} \quad ((F_1 \wedge F_2) \vee (\neg F_1 \wedge \neg F_2))$$

$$\left(\bigvee_{i=1}^n F_i \right) \quad \text{statt} \quad (\dots ((F_1 \vee F_2) \vee F_3) \vee \dots \vee F_n)$$

$$\left(\bigwedge_{i=1}^n F_i \right) \quad \text{statt} \quad (\dots ((F_1 \wedge F_2) \wedge F_3) \wedge \dots \wedge F_n)$$

Bisher sind Formeln lediglich Zeichenreihen:

- Sie haben keinen Inhalt, keine Interpretation.
- Wir müssen erst noch die Bedeutung festlegen.

Die Elemente der Menge $\{0, 1\}$ heißen Wahrheitswerte.

Eine *Belegung* ist eine Funktion $\mathcal{A} : \mathbb{D} \rightarrow \{0, 1\}$, die den atomaren Formeln $\mathbb{D}$ einen Wahrheitswert 0 oder 1 zuweist.

Wir müssen die Funktion $\mathcal{A}$ so erweitern, dass sie auch für Formeln und Teilformeln gilt:

- $\mathcal{A}((F \wedge G)) = \begin{cases} 1, & \text{falls } \mathcal{A}(F) = 1 \text{ und } \mathcal{A}(G) = 1 \\ 0, & \text{sonst} \end{cases}$
- $\mathcal{A}((F \vee G)) = \begin{cases} 1, & \text{falls } \mathcal{A}(F) = 1 \text{ oder } \mathcal{A}(G) = 1 \\ 0, & \text{sonst} \end{cases}$
- $\mathcal{A}(\neg F) = \begin{cases} 1, & \text{falls } \mathcal{A}(F) = 0 \\ 0, & \text{sonst} \end{cases}$

Beispiel: Es sei $\mathcal{A}(A) = 1$, $\mathcal{A}(B) = 1$ und $\mathcal{A}(C) = 0$, dann gilt:

$$\begin{aligned} & \mathcal{A}(\neg((A \wedge B) \vee C)) \\ &= \begin{cases} 1, & \text{falls } \mathcal{A}(((A \wedge B) \vee C)) = 0 \\ 0, & \text{sonst} \end{cases} \\ &= \begin{cases} 0, & \text{falls } \mathcal{A}(((A \wedge B) \vee C)) = 1 \\ 1, & \text{sonst} \end{cases} \\ &= \begin{cases} 0, & \text{falls } \mathcal{A}((A \wedge B)) = 1 \text{ oder } \mathcal{A}(C) = 1 \\ 1, & \text{sonst} \end{cases} \\ &= \begin{cases} 0, & \text{falls } \mathcal{A}((A \wedge B)) = 1 \text{ (da } \mathcal{A}(C) = 0) \\ 1, & \text{sonst} \end{cases} \\ &= \begin{cases} 0, & \text{falls } \mathcal{A}(A) = 1 \text{ und } \mathcal{A}(B) = 1 \\ 1, & \text{sonst} \end{cases} \\ &= 0 \end{aligned}$$

Wir können die Wirkung der Operatoren $\wedge$, $\vee$, $\neg$ auch durch Verknüpfungstafeln darstellen:

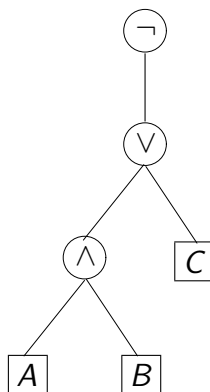
$\mathcal{A}(X)$	$\mathcal{A}(Y)$	$\mathcal{A}(\neg X)$	$\mathcal{A}(X \wedge Y)$	$\mathcal{A}(X \vee Y)$
0	0	1	0	0
0	1	1	0	1
1	0	0	0	1
1	1	0	1	1

Ebenso die Wirkung der Operatoren $\rightarrow$ und $\leftrightarrow$:

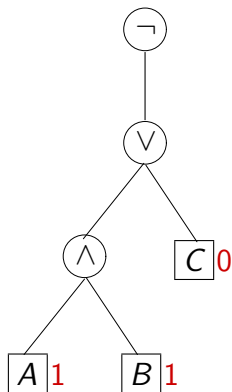
$\mathcal{A}(X)$	$\mathcal{A}(Y)$	$\mathcal{A}(X \rightarrow Y)$	$\mathcal{A}(X \leftrightarrow Y)$
0	0	1	1
0	1	1	0
1	0	0	0
1	1	1	1

Beachte: Aus einer falschen Aussage kann man alles folgern!

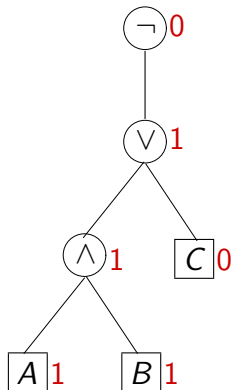
Man kann die Art und Weise, wie Formeln aus einfacheren Teilformeln aufgebaut sind, auch durch eine Baumstruktur darstellen:



Um den Wert der Formel bei einer gegebenen Belegung zu bestimmen, markiere zunächst die Blätter mit den durch die Belegung gegebenen Wahrheitswerten.



Markiere dann von unten nach oben alle Knoten anhand obiger Verknüpfungstafeln.



Definition: Sei F eine Formel und $\mathcal{A}$ eine Belegung.

- Falls $\mathcal{A}$ für alle in F vorkommenden atomaren Formeln definiert ist, heißt $\mathcal{A}$ zu F *passend*.
- Ist $\mathcal{A}$ zu F passend und gilt $\mathcal{A}(F) = 1$, so schreiben wir auch $\mathcal{A} \models F$ und sagen *$\mathcal{A}$ ist ein Modell für F* oder *$\mathcal{A}$ erfüllt F* .
- Falls $\mathcal{A}(F) = 0$, so schreiben wir $\mathcal{A} \not\models F$ und sagen *$\mathcal{A}$ ist kein Modell für F* oder *$\mathcal{A}$ erfüllt F nicht*.

Eine Formel F heißt

- *erfüllbar*, falls F mindestens ein Modell besitzt.
- *unerfüllbar*, falls es kein Modell für F gibt.
- *allgemeingültig* oder *Tautologie*, falls jede Belegung ein Modell für F ist. Wir schreiben dann: $\models F$
- *falsifizierbar*, falls es eine Belegung gibt, die kein Modell für die Formel ist.

Übung: Falls eine Formel F aus n atomaren Formeln aufgebaut ist, wieviele verschiedene Belegungen gibt es für F ?

Übung: Wie viele verschiedene Formeln mit den atomaren Formeln $A_1, \dots, A_n$ und verschiedenen Wahrheitswertverläufen gibt es?

Übung: zeigen Sie folgende wichtige Tautologien

Modus Ponens $(F \wedge (F \rightarrow G)) \rightarrow G$

Modus Tollens $((F \rightarrow G) \wedge \neg G) \rightarrow \neg F$

Und-Elimination $(F \wedge G) \rightarrow F$

Oder-Introduktion $F \rightarrow (F \vee G)$

Resolutionsregel $((F \rightarrow G) \wedge (\neg F \rightarrow H)) \rightarrow (G \vee H)$

Da wir in wissensbasierten Systemen nicht nur eine, sondern viele Regeln haben, müssen wir obige Begriffe auf Formelmengen erweitern. Sei $\mathcal{F}$ eine Menge von Formeln.

- Dann ist $\mathcal{A}$ ein *Modell* für $\mathcal{F}$, falls für alle $F \in \mathcal{F}$ gilt $\mathcal{A} \models F$, d.h. $\mathcal{A}$ erfüllt jede Formel in $\mathcal{F}$.

Insbesondere ist dann $\mathcal{A}$ passend für jede Formel $F \in \mathcal{F}$.

- $\mathcal{F}$ heißt *erfüllbar*, falls $\mathcal{F}$ mindestens ein Modell besitzt, andernfalls heißt $\mathcal{F}$ *unerfüllbar*.

Eine Formel G heißt *Folgerung* der Formeln $F_1, \dots, F_k$, falls für jede passende Belegung $\mathcal{A}$ für $F_1, \dots, F_k$ und G gilt:

$$\mathcal{A} \text{ ist Modell für } \{F_1, \dots, F_k\} \Rightarrow \mathcal{A} \text{ ist Modell für } G$$

Übung: Wir wollen uns ein Haustier anschaffen und machen einige Überlegungen:

- Es soll nur ein Hund (H), eine Katze (K) oder ein Hamster (M) sein.
- Besitzer wertvoller Möbel (W) sollten keine Katze anschaffen, da diese die Möbel zerkratzt.
- Ein Hund erfordert ein freistehendes Haus (F), damit kein Nachbar durch das Bellen gestört wird.

Wir vermuten: Für einen Besitzer wertvoller Möbel ohne freistehendes Haus kommt nur ein Hamster in Frage.

- Wie lauten die Aussagen als aussagenlogische Formeln?
- Wie lautet die Hypothese als aussagenlogische Formel?
- Folgt die Hypothese aus der Formelmenge?

Übung: Inspektor Columbo hat drei Personen in Verdacht, die eine Straftat (evtl. gemeinsam) begangen haben. Ihm stehen folgende Informationen zur Verfügung:

- Wenn A schuldig ist und B unschuldig, dann ist C schuldig.
- C arbeitet niemals allein.
- A arbeitet niemals mit C .
- Nur A , B oder C kommen als Täter in Frage.

Inspektor Columbo vermutet, dass B einer der Täter ist. Helfen Sie ihm, indem Sie die Sachverhalte als aussagenlogische Formeln darstellen. Folgt die Hypothese aus der Formelmenge?

Hatte B einen Komplizen? Begründen Sie Ihre Vermutung.

Übung: Geben Sie eine Formelmenge $\mathcal{F}$ an, so dass jede echte Teilmenge von $\mathcal{F}$ erfüllbar ist, $\mathcal{F}$ selbst jedoch nicht.

Übung: Ist folgende unendliche Formelmenge $\mathcal{F}$ erfüllbar?

$$\mathcal{F} = \{A_1 \vee A_2, \neg A_2 \vee \neg A_3, A_3 \vee A_4, \neg A_4 \vee \neg A_5, \dots\}$$

Übung: zeigen Sie, dass folgende Behauptungen äquivalent sind:

- G ist eine Folgerung von $F_1, \dots, F_k$
- $((\bigwedge_{i=1}^k F_i) \rightarrow G)$ ist eine Tautologie
- $((\bigwedge_{i=1}^k F_i) \wedge \neg G)$ ist unerfüllbar

Anmerkungen:

- Wir sehen an unseren Beispielen (Haustier, Inspektor Columbo), dass das Testen von Belegungen ineffizient ist.
- Für maschinelle Inferenz wollen wir Techniken nutzen, die allein auf der Syntax der Formeln beruhen.
- Wir suchen eine Folge von syntaktischen Umformungen, die die Hypothese beweisen.

Daher beschäftigen wir uns im Folgenden mit Äquivalenz und Normalformen.

Aussagenlogik

- Grundbegriffe
- *Äquivalenz und Normalformen*
- Hornformeln
- Resolution

Syntaktisch unterschiedliche Formeln können dasselbe aussagen:

- $(F \wedge G)$ und $(G \wedge F)$ bedeuten dasselbe, obwohl es syntaktisch zwei verschiedene Objekte sind.
- Auch $\neg(F \vee G)$ und $\neg F \wedge \neg G$ sagen dasselbe aus, wie man an den Wahrheitstafeln erkennen kann.

F	G	$\neg(F \vee G)$	$\neg F \wedge \neg G$
0	0	1	1
0	1	0	0
1	0	0	0
1	1	0	0

Im Weiteren: Um eine Hypothese zu beweisen, wollen wir nicht mehr alle Belegungen testen, sondern eine Folge von syntaktischen Umformungen durchführen.

Definition: Zwei Formeln F und G heißen *(semantisch) äquivalent*, falls für alle Belegungen $\mathcal{A}$, die für F und auch für G passend sind, $\mathcal{A}(F) = \mathcal{A}(G)$ gilt. Wir schreiben dann $F \equiv G$.

$$(F \wedge F) \equiv F \quad \text{Idempotenz}$$

$$(F \vee F) \equiv F$$

$$(F \wedge G) \equiv (G \wedge F) \quad \text{Kommutativität}$$

$$(F \vee G) \equiv (G \vee F)$$

$$((F \wedge G) \wedge H) \equiv (F \wedge (G \wedge H)) \quad \text{Assoziativität}$$

$$((F \vee G) \vee H) \equiv (F \vee (G \vee H))$$

$$(F \wedge (F \vee G)) \equiv F \quad \text{Absorption}$$

$$(F \vee (F \wedge G)) \equiv F$$

$$\begin{aligned}(F \wedge (G \vee H)) &\equiv ((F \wedge G) \vee (F \wedge H)) && \text{Distributivität} \\ (F \vee (G \wedge H)) &\equiv ((F \vee G) \wedge (F \vee H))\end{aligned}$$

$$\neg\neg F \equiv F \quad \text{Doppelnegation}$$

$$\begin{aligned}\neg(F \wedge G) &\equiv (\neg F \vee \neg G) \\ \neg(F \vee G) &\equiv (\neg F \wedge \neg G)\end{aligned} \quad \begin{array}{l} \text{deMorgansche} \\ \text{Regeln} \end{array}$$

$$\begin{aligned}(F \vee G) &\equiv F, \text{ falls } F \text{ Tautologie} \\ (F \wedge G) &\equiv G, \text{ falls } F \text{ Tautologie}\end{aligned} \quad \text{Tautologieregeln}$$

$$\begin{aligned}(F \vee G) &\equiv G, \text{ falls } F \text{ unerfüllbar} \\ (F \wedge G) &\equiv F, \text{ falls } F \text{ unerfüllbar}\end{aligned} \quad \begin{array}{l} \text{Unerfüllbarkeits-} \\ \text{regeln} \end{array}$$

Beweis: Alle obigen Äquivalenzen können leicht mittels Wahrheitstafeln nachgeprüft werden.

Beispiel Absorption: $F \wedge (F \vee G) \equiv F$

$\mathcal{A}(F)$	$\mathcal{A}(G)$	$\mathcal{A}((F \vee G))$	$\mathcal{A}((F \wedge (F \vee G)))$
0	0	0	0
0	1	1	0
1	0	1	1
1	1	1	1

Da die erste und die letzte Spalte übereinstimmen, ist die Äquivalenz der beiden Formeln gezeigt.

Übung: Zeigen Sie sowohl durch Wahrheitstabeln als auch durch Anwendung obiger Umformungsregeln, dass

$$((A \vee \neg(B \wedge A)) \wedge (C \vee (D \vee C)))$$

äquivalent ist zu $(C \vee D)$.

Übung: Zeigen Sie, dass es zu jeder Formel F eine äquivalente Formel G gibt, die nur die Operatoren $\neg$ und $\rightarrow$ enthält.

Übung: Zeigen Sie sowohl durch Wahrheitstabeln als auch durch Anwendung obiger Umformungsregeln, dass die beiden folgenden Formeln äquivalent sind.

- $(\neg P \vee Q) \leftrightarrow R$
- $(\neg P \wedge R) \vee (Q \wedge R) \vee (P \wedge \neg Q \wedge \neg R)$

Übung: Beweisen Sie die folgenden Verallgemeinerungen der deMorganschen Gesetze:

$$\neg \left(\bigvee_{i=1}^n F_i \right) \equiv \left(\bigwedge_{i=1}^n \neg F_i \right)$$

$$\neg \left(\bigwedge_{i=1}^n F_i \right) \equiv \left(\bigvee_{i=1}^n \neg F_i \right)$$

Übung: Beweisen Sie die folgenden Verallgemeinerungen der Distributivgesetze:

$$\left(\left(\bigvee_{i=1}^m F_i \right) \wedge \left(\bigvee_{j=1}^n G_j \right) \right) \equiv \left(\bigvee_{i=1}^m \left(\bigvee_{j=1}^n (F_i \wedge G_j) \right) \right)$$

$$\left(\left(\bigwedge_{i=1}^m F_i \right) \vee \left(\bigwedge_{j=1}^n G_j \right) \right) \equiv \left(\bigwedge_{i=1}^m \left(\bigwedge_{j=1}^n (F_i \vee G_j) \right) \right)$$

Im Folgenden: Jede, auch noch so kompliziert aussehende Formel kann in eine gewisse Normalform überführt werden. Obige Umformungsregeln sind dafür ausreichend.

- Ein *Literal* ist eine atomare Formel oder die Negation einer atomaren Formel.
- Eine Formel F ist in *konjunktiver Normalform* (KNF), falls sie eine Konjunktion von Disjunktionen von Literalen ist:

$$F = \left(\bigwedge_{i=1}^n \left(\bigvee_{j=1}^{m_i} L_{i,j} \right) \right)$$

- Eine Formel F ist in *disjunktiver Normalform* (DNF), falls sie eine Disjunktion von Konjunktionen von Literalen ist:

$$F = \left(\bigvee_{i=1}^n \left(\bigwedge_{j=1}^{m_i} L_{i,j} \right) \right)$$

- Die $L_{i,j}$ sind Literale mit $L_{i,j} \in \{A_1, A_2, \dots\} \cup \{\neg A_1, \neg A_2, \dots\}$.

Satz: Für jede Formel F gibt es eine äquivalente Formel in KNF und eine äquivalente Formel in DNF. (Beweis durch Induktion über den Formelaufbau von F , wird hier nicht dargestellt.)

Umformungsmethode zur Herstellung der KNF:

- Ersetze in F jedes Vorkommen von

$$\neg\neg G \quad \text{durch} \quad G$$

$$\neg(G \wedge H) \quad \text{durch} \quad (\neg G \vee \neg H)$$

$$\neg(G \vee H) \quad \text{durch} \quad (\neg G \wedge \neg H)$$

bis keine derartige Teilformel mehr vorkommt.

- Ersetze in F jedes Vorkommen von

$$(F \vee (G \wedge H)) \quad \text{durch} \quad ((F \vee G) \wedge (F \vee H))$$

$$((F \wedge G) \vee H) \quad \text{durch} \quad ((F \vee H) \wedge (G \vee H))$$

bis keine derartige Teilformel mehr vorkommt.

Oder wir erstellen die Wahrheitstafel zur Formel F und lesen direkt die Teilformeln der disjunktiven Normalform ab:

A	B	C	F	
0	0	0	1	$\rightarrow \neg A \wedge \neg B \wedge \neg C$
0	0	1	0	
0	1	0	0	
0	1	1	1	$\rightarrow \neg A \wedge B \wedge C$
1	0	0	1	$\rightarrow A \wedge \neg B \wedge \neg C$
1	0	1	1	$\rightarrow A \wedge \neg B \wedge C$
1	1	0	0	
1	1	1	0	

Die abgelesenen Teilformeln werden mittels $\vee$ verknüpft:

$$(\neg A \wedge \neg B \wedge \neg C) \vee (\neg A \wedge B \wedge C) \vee (A \wedge \neg B \wedge \neg C) \vee (A \wedge \neg B \wedge C)$$

Die Korrektheit dieses Verfahrens ist intuitiv klar.

Analog können wir aus der Wahrheitstafel einer Formel F direkt die Teilformeln der konjunktiven Normalform ablesen:

A	B	C	F	
0	0	0	1	
0	0	1	0	$\rightarrow A \vee B \vee \neg C$
0	1	0	0	$\rightarrow A \vee \neg B \vee C$
0	1	1	1	
1	0	0	1	
1	0	1	1	
1	1	0	0	$\rightarrow \neg A \vee \neg B \vee C$
1	1	1	0	$\rightarrow \neg A \vee \neg B \vee \neg C$

Die abgelesenen Teilformeln werden mittels $\wedge$ verknüpft:

$$(A \vee B \vee \neg C) \wedge (A \vee \neg B \vee C) \wedge (\neg A \vee \neg B \vee C) \wedge (\neg A \vee \neg B \vee \neg C)$$

Dieses Verfahren ist intuitiv vielleicht nicht sofort klar, aber ...

... wir können hier die Darstellung in DNF benutzen, um uns die Korrektheit des Verfahrens klar zu machen:

<i>A</i>	<i>B</i>	<i>C</i>	<i>F</i>				
0	0	0	1				
0	0	1	0	→	$A \vee B \vee \neg C$	≡	$\neg(\neg A \wedge \neg B \wedge C)$
0	1	0	0	→	$A \vee \neg B \vee C$	≡	$\neg(\neg A \wedge B \wedge \neg C)$
0	1	1	1				
1	0	0	1				
1	0	1	1				
1	1	0	0	→	$\neg A \vee \neg B \vee C$	≡	$\neg(A \wedge B \wedge \neg C)$
1	1	1	0	→	$\neg A \vee \neg B \vee \neg C$	≡	$\neg(A \wedge B \wedge C)$

Wir erstellen die KNF für die negierte Formel und erhalten:

$$\neg((\neg A \wedge \neg B \wedge C) \vee (\neg A \wedge B \wedge \neg C) \vee (A \wedge B \wedge \neg C) \vee (A \wedge B \wedge C))$$

Wenn wir jetzt die Negation nach den deMorganschen Regeln in die Klammer ziehen, erhalten wir genau das gewünschte.

Anmerkung: Die so erzeugten KNF- bzw. DNF-Formeln sind nicht notwendigerweise die kürzesten. Z.B. kann man

$$(A \vee B \vee \neg C) \wedge (A \vee \neg B \vee C) \wedge (\neg A \vee \neg B \vee C) \wedge (\neg A \vee \neg B \vee \neg C)$$

vereinfachen zu

$$(A \vee B \vee \neg C) \wedge (A \vee \neg B \vee C) \wedge (\neg A \vee \neg B)$$

Übung: Erzeugen Sie mittels Umformung und mittels einer Wahrheitstafel zu der Formel

$$F = ((\neg A \rightarrow B) \wedge ((A \wedge \neg C) \leftrightarrow B))$$

eine äquivalente DNF und KNF.

Aussagenlogik

- Grundbegriffe
- Äquivalenz und Normalformen
- *Hornformeln*
- Resolution

Definition: Eine Formel F ist eine *Hornformel*, falls F in KNF ist, und jedes Disjunktionsglied in F höchstens ein positives Literal enthält. Wurde benannt nach dem Logiker Alfred Horn.

Beispiel:

$$F = (A \vee \neg B) \wedge (\neg C \vee \neg A \vee D) \wedge (\neg A \vee \neg B) \wedge D \wedge \neg E$$

Hornformeln können anschaulicher geschrieben werden:

$$F = (B \rightarrow A) \wedge (C \wedge A \rightarrow D) \wedge (A \wedge B \rightarrow 0) \wedge (1 \rightarrow D) \wedge (E \rightarrow 0)$$

Dabei steht 0 für eine beliebige unerfüllbare Formel und 1 für eine beliebige Tautologie.

Algorithmischer Test für die Erfüllbarkeit oder Unerfüllbarkeit einer gegebenen Hornformel F :

initial: für alle Teilformeln $1 \rightarrow A_i$ aus F markiere A_i
A_i muss mit 1 belegt sein, um F zu erfüllen

wiederhole

- markiere y , falls es eine Teilformel $(A_{i_1} \wedge \dots \wedge A_{i_k}) \rightarrow y$ gibt, bei der $A_{i_1}, \dots, A_{i_k}$ bereits markiert sind
y muss mit 1 belegt sein, um F zu erfüllen
- falls es eine Teilformel $(A_{i_1} \wedge \dots \wedge A_{i_k}) \rightarrow 0$ gibt, bei der $A_{i_1}, \dots, A_{i_k}$ bereits markiert sind, gib *unerfüllbar* aus und stoppe

gib *erfüllbar* aus und stoppe

Die Korrektheit des Algorithmus ergibt sich anhand der obigen Bemerkungen.

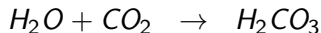
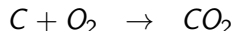
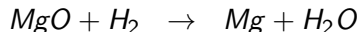
Übung: Wenden Sie den Markierungsalgorithmus auf die Formel

$$(\neg A \vee \neg B \vee \neg D) \wedge \neg E \wedge (\neg C \vee A) \wedge C \wedge B \wedge (\neg G \vee D) \wedge G$$

an. Die Wahrheitstafel hätte für diese Formel $2^6 = 64$ Zeilen!

Übung: Geben Sie eine Formel an, zu der es keine äquivalente Hornformel gibt und begründen Sie, warum das so ist.

Übung: Angenommen, wir könnten die folgenden chemischen Reaktionen durchführen:



Ist es dann möglich, H_2CO_3 aus den Grundstoffen MgO , H_2 , O_2 und C herzustellen?

Aussagenlogik

- Grundbegriffe
- Äquivalenz und Normalformen
- Hornformeln
- *Resolution*

Ziel unserer Bemühungen in diesem ganzen Kapitel ist: Wir wollen die Unerfüllbarkeit einer gegebenen Formelmenge nachweisen!

Warum können wir uns auf Unerfüllbarkeitstests beschränken?

- Um zu testen, ob eine Formel F eine Tautologie ist, testen wir einfach $\neg F$ auf Unerfüllbarkeit.
- Um zu testen, ob eine Formel G eine Folgerung einer Formelmenge $\mathcal{F} = \{F_1, F_2, \dots, F_k\}$ ist, testen wir die Unerfüllbarkeit von $F_1 \wedge F_2 \wedge \dots \wedge F_k \wedge \neg G$.

- Resolution ist eine einfach anzuwendende syntaktische Umformungsregel.
- Generiere aus zwei Formeln eine dritte Formel, die dann wieder als Eingabe in weitere Resolutionsschritte dienen kann.
- Voraussetzung für die Anwendung der Resolution: Die Formeln müssen in konjunktiver Normalform vorliegen.
- Eine Kollektion rein mechanisch anzuwendender syntaktischer Umformungsregeln nennen wir *Kalkül*.

Ein Kalkül ergänzt die Aussagenlogik oder die Prädikatenlogik um den syntaktischen Begriff des Ableitens, der dem semantischen Folgerungsbegriff entspricht.

- Die Begriffe „wahr“ und „falsch“ kommen in diesem System zunächst gar nicht vor.
- Stattdessen werden Axiome gesetzt, die einfach als Zeichenketten angesehen werden.
- Aus diesen Zeichenketten können weitere ableitbare Zeichenketten aufgrund von bestimmten Schlussregeln (Inferenzregeln) hergeleitet werden.

- In dem formalen System sollen nur Zeichenketten (Sätze) hergeleitet werden können, die bei einer plausiblen Interpretation auch wahr sind. → *Korrektheit*
- Alle Sätze, die als „wahr“ interpretierbar sind, sollen auch hergeleitet werden können. → *Vollständigkeit*

Klassische Aussagenlogik:

- Die Symbole $\wedge$, $\vee$ und $\neg$ mit den bekannten Interpretationen bilden die Grundlage eines Kalküls (nach George Boole).
- Kalküle sind korrekt als auch vollständig.

Komplexere logische Systeme wie bspw. die Mengenlehre:

- Es ist unmöglich, einen vollständigen Kalkül aufzustellen, der auch korrekt ist,
- siehe Gödelscher Unvollständigkeitssatz.

Die Definition eines Kalküls macht nur dann Sinn, wenn dessen Korrektheit und Vollständigkeit in Bezug zur Aufgabenstellung nachgewiesen werden kann:

- *Aufgabenstellung*: Wir wollen die Unerfüllbarkeit einer gegebenen Formelmenge nachweisen.
- *Korrektheit*: Keine erfüllbare Formelmenge wird durch den Kalkül als vermeintlich unerfüllbar nachgewiesen.
- *Vollständigkeit*: Jede unerfüllbare Formelmenge wird durch den Kalkül als solche nachgewiesen.

Seien $A_1, A_2, A_3, \dots$ aussagenlogische Variablen. Eine endliche Menge $K \subset \{A_1, A_2, A_3, \dots\} \cup \{\neg A_1, \neg A_2, \neg A_3, \dots\}$ heißt eine *Klausel*.

Die Disjunktionsglieder einer konjunktiven Normalform können über Klauseln dargestellt werden.

Beispiel: Die Formel $(A \vee B) \wedge (\neg A \vee C)$ führt zu den Klauseln $\{A, B\}$ und $\{\neg A, C\}$.

Die Menge der Klauseln zu den Disjunktionsgliedern heißt *Klauselmenge der aussagenlogischen Formel*.

Beispiel: Die Formel $(A \vee B) \wedge (\neg A \vee C)$ führt zur Klauselmenge $\{\{A, B\}, \{\neg A, C\}\}$.

Seien K_1 , K_2 und R Klauseln. R heißt *Resolvente* von K_1 und K_2 , falls es ein Literal L gibt mit $L \in K_1$ und $\bar{L} \in K_2$ und R die Form hat:

$$R = (K_1 - \{L\}) \cup (K_2 - \{\bar{L}\})$$

Beispiel:

$$\begin{array}{ccc} \{A, \neg B, C\} & & \{B, C, D\} \\ & \searrow \quad \swarrow & \\ & \{A, C, D\} & \end{array}$$

Übung: Kann bei der Resolution zweier Hornformeln eine Formel entstehen, die keine Hornformel ist? Begründen Sie Ihre Meinung.

Sei $\mathcal{K}$ eine Klauselmenge. Dann definieren wir:

$$\text{Res}(\mathcal{K}) := \mathcal{K} \cup \{R : R \text{ ist Resolvente zweier Klauseln von } \mathcal{K}\}$$

Außerdem setzen wir:

$$\text{Res}^0(\mathcal{K}) := \mathcal{K}$$

$$\text{Res}^{n+1}(\mathcal{K}) := \text{Res}(\text{Res}^n(\mathcal{K}))$$

$$\text{Res}^*(\mathcal{K}) := \bigcup_{n=1}^{\infty} \text{Res}^n(\mathcal{K})$$

Übung: Bestimmen Sie für folgende Klauselmenge $\mathcal{K}$ die Mengen $\text{Res}^n(\mathcal{K})$ für $n = 0, 1, 2$.

$$\mathcal{K} = \{\{A, \neg B, C\}, \{B, C\}, \{\neg A, C\}, \{B, \neg C\}, \{\neg C\}\}$$

Übung: Geben Sie sämtliche Resolventen an, die aus den Klauseln der folgenden Klauselmenge gewonnen werden können:

$$\{\{A, B, \neg E\}, \{\neg A, \neg B, \neg D\}, \{A, D\}\}$$

Übung: Beweisen Sie, dass es für jede endliche Klauselmenge $\mathcal{K}$ ein $k \geq 0$ gibt mit

$$\text{Res}^k(\mathcal{K}) = \text{Res}^{k+1}(\mathcal{K}) = \dots = \text{Res}^*(\mathcal{K}).$$

Sei F eine Formel in konjunktiver Normalform, dargestellt als Klauselmenge. Ferner sei R eine Resolvente zweier Klauseln K_1 und K_2 aus F . Dann sind F und $F \cup \{R\}$ äquivalent.

Beweis:

- Falls F unerfüllbar ist: Wenn wir F mit einer weiteren Klausel $\wedge$ -verknüpfen, ist auch die neue Formel unerfüllbar.
- Falls F mit einer Variablenbelegung $\mathcal{A}$ erfüllbar ist:
 - Beide Eingangsklauseln K_1 und K_2 sind mit $\mathcal{A}$ erfüllt.
 - Das für die Resolution benutzte Literal L kommt einmal negiert und einmal nicht-negiert in K_1 bzw. K_2 vor, also ist entweder $\mathcal{A}(L) = 0$ oder $\mathcal{A}(\bar{L}) = 0$.
 - Die restliche Klausel beim nicht-erfüllten Auftreten muss erfüllt sein, damit die Klausel und damit auch F erfüllt ist.
 - Auch die Resolvente ist erfüllt, da in jedem Disjunktionsglied nur ein einziges Literal erfüllt sein muss.

Satz: Eine Formel in konjunktiver Normalform mit Klauselmenge $\mathcal{K}$ ist genau dann unerfüllbar, falls mittels Resolution die leere Menge abgeleitet werden kann, also $\emptyset \in \text{Res}^*(\mathcal{K})$.

Wir zeigen zunächst: $\mathcal{K}$ unerfüllbar $\Leftrightarrow \emptyset \in \text{Res}^*(\mathcal{K})$

Sei also $\emptyset \in \text{Res}^k(\mathcal{K})$ und k minimal.

- In $\text{Res}^{k-1}(\mathcal{K})$ gibt es zwei Klauseln vom Typ $\{A\}$ und $\{\neg A\}$.
- $A \wedge \neg A$ ist unerfüllbar, also ist $\text{Res}^{k-1}(\mathcal{K})$ unerfüllbar.
- Aufgrund des Resolutions-Lemmas gilt:

$$\mathcal{K} \equiv \text{Res}^1(\mathcal{K}) \equiv \text{Res}^2(\mathcal{K}) \equiv \dots \equiv \text{Res}^n(\mathcal{K}) \equiv \dots$$

- Somit ist auch $\mathcal{K}$ unerfüllbar.

Noch zu zeigen: $\mathcal{K}$ unerfüllbar $\Rightarrow \emptyset \in \text{Res}^*(\mathcal{K})$

Beweis durch Induktion über die Anzahl der Variablen n einer Klauselmeng $\mathcal{K}$.

- Induktionsanfang für $n = 1$:

$\mathcal{K} = \{\{A\}, \{\neg A\}\}$ ist die einzige unerfüllbare Klauselmeng (falls die Variable A genannt wird) und $\emptyset$ entsteht bei der Resolution.

- Induktionsannahme:

Eine fortgesetzte Resolution liefert für jede unerfüllbare Klauselmeng mit n Variablen $A_1, A_2, \dots, A_n$ u.a. die leere Meng.

- Induktionsschritt: Zeige dies für $n + 1$ Variablen.

Induktionsschritt: Sei $\mathcal{K}$ eine unerfüllbare Klauselmenge mit $n + 1$ Variablen $A_1, A_2, \dots, A_{n+1}$.

- $\mathcal{K}_0$ entstehe aus $\mathcal{K}$ durch Weglassen von A_{n+1} aus allen Klauseln und Weglassen aller Klauseln mit $\neg A_{n+1}$.
→ belege A_{n+1} mit 0
- $\mathcal{K}_1$ entstehe aus $\mathcal{K}$ durch Weglassen von $\neg A_{n+1}$ aus allen Klauseln und Weglassen aller Klauseln mit A_{n+1} .
→ belege A_{n+1} mit 1

Gegeben sei eine beliebige Belegung $\mathcal{A}$ der Variablen. Da $\mathcal{K}$ unerfüllbar ist, gilt:

- $\mathcal{K}_0$ ist unerfüllbar, falls $\mathcal{A}(A_{n+1}) = 0$.
- $\mathcal{K}_1$ ist unerfüllbar, falls $\mathcal{A}(A_{n+1}) = 1$.

Nach Induktionsannahme kann man also in beiden Fällen mittels Resolution $\emptyset$ erzeugen:

$$\emptyset \in \text{Res}^*(\mathcal{K}_0) \quad \text{und} \quad \emptyset \in \text{Res}^*(\mathcal{K}_1)$$

Fügt man A_{n+1} und $\neg A_{n+1}$ wieder hinzu, so liefert die Resolution für das erweiterte $\mathcal{K}_0$ eine Menge, in der $\{A_{n+1}\}$ enthalten ist, für das erweiterte $\mathcal{K}_1$ ist $\{\neg A_{n+1}\}$ enthalten. Also

$$\{A_{n+1}\} \in \text{Res}^*(\mathcal{K}_0) \quad \text{und} \quad \{\neg A_{n+1}\} \in \text{Res}^*(\mathcal{K}_1)$$

Die Resolvente von beiden Mengen ist $\emptyset$, also $\emptyset \in \text{Res}^*(\mathcal{K})$.

Unerfüllbarkeit einer Klauselmenge: Zeigen Sie mittels Resolution die Unerfüllbarkeit der Klauselmenge

$$\{A, B, \neg C\}, \{\neg A\}, \{A, B, C\}, \{A, \neg B\},$$

also der Formel

$$(A \vee B \vee \neg C) \wedge (\neg A) \wedge (A \vee B \vee C) \wedge (A \vee \neg B).$$

Haustierbeispiel: Leiten Sie die Hypothese „für einen Besitzer wertvoller Möbel ohne freistehendes Haus kommt nur ein Hamster in Frage“ mittels Resolution her.

Diätplan: „Das ist ja unmöglich, mit dieser Baukasten-Diät“, seufzt Robert:

- Wenn ich einen Apfel esse, muss ich auch einen Salzhering nehmen. Wegen der Mineralien.
- Esse ich eine Banane, muss ich eine Scheibe Knäckebrot hinterher essen - oder ich lasse den Salzhering liegen.
- Aber wenn ich ein Tortenstück esse, darf ich kein Knäckebrot nehmen - zu viele Kohlehydrate.
- Esse ich einen Apfel, soll ich auch ein Tortenstück nehmen.
- Wenn ich keinen Milchreis esse, soll ich eine Banane nehmen.
- Ich esse auf jeden Fall einen Apfel, aber keinen Milchreis - den vertrage ich nicht.

„Pack den Diätplan weg“, sagt Sabine, „er ist für dich wirklich unmöglich!“. Frage: Hat Sabine Recht?

Wir können keine Aussagen über ganze Klassen von Objekten machen, so dass Schlussfolgerungen für individuelle Objekte möglich sind.

Beispiel: Gegeben seien folgende Fakten bzw. Regeln:

- Martin ist Informatiker.
- Peter ist Informatiker.
- Jeder Informatiker kann programmieren.

Folgerungen:

- Martin kann programmieren.
- Peter kann programmieren.

- Geschichte und Anwendungen
- Zustandsraumsuche
- *Aussagen- und Prädikatenlogik*
- Einführung in Prolog

Prädikatenlogik

- *Grundbegriffe*
- Normalformen
- Unentscheidbarkeit
- Herbrand-Theorie
- Resolution

Die Prädikatenlogik ist eine Erweiterung der Aussagenlogik. Hinzu kommen der Allquantor $\forall$, der Existenzquantor $\exists$ sowie Funktions- und Prädikatsymbole.

Analog zur Aussagenlogik kann auch für die Prädikatenlogik deren Syntax und Semantik mathematisch präzise definiert werden.

Syntax der Prädikatenlogik:

- Eine *Variable* hat die Form x_i mit $i = 1, 2, 3, \dots$
- Ein *Prädikatsymbol* hat die Form P_i^k mit $i = 1, 2, 3, \dots$ und $k = 0, 1, 2, \dots$
- Ein *Funktionssymbol* hat die Form f_i^k mit $i = 1, 2, 3, \dots$ und $k = 0, 1, 2, \dots$
- Man nennt i den *Unterscheidungsindex* und k die *Stellenzahl* oder auch *Stelligkeit*.

Der Begriff *Term* wird durch einen induktiven Prozess definiert:

- Jede Variable ist ein Term.
- Falls f ein Funktionssymbol ist mit Stellenzahl k , und falls $t_1, \dots, t_k$ Terme sind, so ist auch $f(t_1, \dots, t_k)$ ein Term.

Auch Funktionssymbole der Stellenzahl 0 sollen eingeschlossen sein, in einem solchen Fall fallen die Klammern weg.

Nullstellige Funktionssymbole heißen auch *Konstanten*.

Auch die Formeln der Prädikatenlogik können wir induktiv definieren:

- Falls P ein Prädikatsymbol der Stelligkeit k ist, und falls $t_1, \dots, t_k$ Terme sind, dann ist $P(t_1, \dots, t_k)$ eine Formel. Solche Formeln nennen wir auch *atomare Formeln*.
- Für jede Formel F ist auch $\neg F$ eine Formel.
- Für alle Formeln F und G sind auch $(F \wedge G)$ und $(F \vee G)$ Formeln.
- Falls x eine Variable ist und F eine Formel, so sind auch $\exists x F$ und $\forall x F$ Formeln.

Falls F eine Formel ist und F als Teil einer Formel G auftritt, so heißt F *Teilformel* von G .

Alle Vorkommen von Variablen in einer Formel werden in *freie* und *gebundene* Vorkommen unterteilt:

- Ein Vorkommen der Variablen x in der Formel F heißt *gebunden*, falls x in einer Teilformel der Form $\exists x G$ oder $\forall x G$ vorkommt,
 - andernfalls heißt dieses Vorkommen von x *frei*.
- Dieselbe Variable kann also in einer Formel an verschiedenen Stellen sowohl frei als auch gebunden vorkommen.

Eine Formel ohne Vorkommen einer freien Variablen heißt *geschlossen* oder auch eine *Aussage*.

Beispiel:

$$F = (\exists x_1 P_5^2(x_1, f_2^1(x_2)) \vee \neg \forall x_2 P_4^2(x_2, f_7^2(f_4^0, f_5^1(x_3))))$$

ist eine Formel. Sämtliche Teilformeln von F sind:

- F
- $\neg \forall x_2 P_4^2(x_2, f_7^2(f_4^0, f_5^1(x_3)))$
- $\exists x_1 P_5^2(x_1, f_2^1(x_2))$
- $\forall x_2 P_4^2(x_2, f_7^2(f_4^0, f_5^1(x_3)))$
- $P_5^2(x_1, f_2^1(x_2))$
- $P_4^2(x_2, f_7^2(f_4^0, f_5^1(x_3)))$

Weiterhin gilt:

- Alle Vorkommen von x_1 in F sind gebunden.
 - Erstes Vorkommen von x_2 ist frei, zweites ist gebunden.
 - x_3 kommt in F frei vor.
- Die Formel F ist nicht geschlossen bzw. keine Aussage.

Alle in F vorkommenden Terme sind:

- x_1
- x_2
- $f_2^1(x_2)$
- $f_7^2(f_4^0, f_5^1(x_3))$
- f_4^0
- $f_5^1(x_3)$
- x_3

Der Term f_4^0 stellt eine Konstante dar.

Wir vereinbaren wieder die vereinfachenden Schreibweisen wie in der Aussagenlogik. Außerdem legen wir fest:

- u, v, w, x, y, z stehen für Variablen.
- a, b, c stehen für Konstanten.
- f, g, h stehen für Funktionssymbole, die Stelligkeit geht immer aus dem Kontext hervor.
- P, Q, R stehen für Prädikatsymbole, auch hier geht die Stelligkeit aus dem Kontext hervor.

Übung: Geben Sie sämtliche Teilformeln und Terme an, die in der folgenden Formel enthalten sind:

$$F = ((Q(x) \vee \exists x \forall y (P(f(x), z) \wedge Q(a))) \vee \forall z R(x, z, g(x)))$$

Welche Teilformeln sind Aussagen? Bestimmen Sie für jedes Vorkommen einer Variablen, ob die Variable frei oder gebunden ist.

Semantik der Prädikatenlogik: Eine *Struktur* ist ein Paar $\mathcal{A} = (U_{\mathcal{A}}, I_{\mathcal{A}})$, wobei $U_{\mathcal{A}}$ eine beliebige aber nicht leere Menge ist, die die *Grundmenge* von $\mathcal{A}$, der *Individuenbereich* oder auch *Universum* genannt wird.

Ferner ist $I_{\mathcal{A}}$ eine Abbildung, die

- jedem im Definitionsbereich von $I_{\mathcal{A}}$ liegenden k -stelligen Prädikatsymbol P ein k -stelliges Prädikat über $U_{\mathcal{A}}$ zuordnet,
- jedem im Definitionsbereich von $I_{\mathcal{A}}$ liegenden k -stelligen Funktionssymbol f eine k -stellige Funktion auf $U_{\mathcal{A}}$ zuordnet,
- jeder Variablen x (sofern $I_{\mathcal{A}}$ auf x definiert ist) ein Element der Grundmenge $U_{\mathcal{A}}$ zuordnet.

Wir schreiben abkürzend statt $I_{\mathcal{A}}(P)$ einfach $P^{\mathcal{A}}$, statt $I_{\mathcal{A}}(f)$ einfach $f^{\mathcal{A}}$ und statt $I_{\mathcal{A}}(x)$ einfach $x^{\mathcal{A}}$.

Sei F eine Formel und $\mathcal{A} = (U_{\mathcal{A}}, I_{\mathcal{A}})$ eine Struktur. $\mathcal{A}$ heißt zu F *passend*, falls $I_{\mathcal{A}}$ für alle in F vorkommenden Prädikatsymbole, Funktionssymbole und freien Variablen definiert ist.

Beispiel: $F = \forall x P(x, f(x)) \wedge Q(g(a, z))$ ist eine Formel. Eine zu F passende Struktur ist bspw. $\mathcal{A} = (U_{\mathcal{A}}, I_{\mathcal{A}})$ mit

$$\begin{aligned}U_{\mathcal{A}} &= \{0, 1, 2, 3, \dots\} = \mathbb{N} \\I_{\mathcal{A}}(P) &= P^{\mathcal{A}} = \{(m, n) \mid m, n \in U_{\mathcal{A}} \text{ und } m < n\} \\I_{\mathcal{A}}(Q) &= Q^{\mathcal{A}} = \{n \in U_{\mathcal{A}} \mid n \text{ ist Primzahl}\} \\I_{\mathcal{A}}(f) &= f^{\mathcal{A}} \text{ mit } f^{\mathcal{A}}(n) := n + 1 \\I_{\mathcal{A}}(g) &= g^{\mathcal{A}} \text{ mit } g^{\mathcal{A}}(m, n) := m + n \\I_{\mathcal{A}}(a) &= a^{\mathcal{A}} = 2 \\I_{\mathcal{A}}(z) &= z^{\mathcal{A}} = 3\end{aligned}$$

Mit anderen Worten: Wir haben Interpretationen der Prädikate und Funktionen angegeben und eine Belegung der freien Variablen und der Konstanten festgelegt. Wir wissen jetzt,

- was die Prädikate P und Q bedeuten,
- wie die Funktionen f und g definiert sind,
- das die Konstante a den Wert 2 hat,
- das die Variable x den Wert 3 hat
- und auf welcher Grundmenge das alles definiert ist.

Intuitiv:

- Jede natürliche Zahl ist kleiner als ihr Nachfolger und die Summe von 2 und 3 ist eine Primzahl.
- In dieser Struktur gilt die Formel F .
- F ist nicht notwendigerweise in jeder Struktur gültig!

Wir wollen den Begriff *gültig* nun formal definieren.

Sei F eine Formel und $\mathcal{A}$ eine zu F passende Struktur. Für jeden Term t aus Variablen und Funktionssymbolen von F definieren wir induktiv den **Wert** von t in der Struktur $\mathcal{A}$, den wir mit $\mathcal{A}(t)$ bezeichnen:

- Falls t eine Variable ist (also $t = x$), so ist $\mathcal{A}(t) = x^{\mathcal{A}}$.
- Falls t die Form $t = f(t_1, \dots, t_k)$ hat, wobei $t_1, \dots, t_k$ Terme und f ein k -stelliges Funktionssymbol ist, so ist

$$\mathcal{A}(t) = f^{\mathcal{A}}(\mathcal{A}(t_1), \dots, \mathcal{A}(t_k)).$$

Dieser Fall schließt auch die Möglichkeit ein, dass f nullstellig ist, also t die Form $t = a$ hat. Dann ist $\mathcal{A}(t) = a^{\mathcal{A}}$.

Analog definieren wir induktiv den (Wahrheits-) **Wert** der Formel F unter der Struktur $\mathcal{A}$, wobei wir ebenfalls die Bezeichnung $\mathcal{A}(F)$ verwenden.

- Falls F die Form $F = P(t_1, \dots, t_k)$ hat mit den Termen $t_1, \dots, t_k$ und einem k -stelligen Prädikatsymbol P , so gilt:

$$\mathcal{A}(F) = \begin{cases} 1, & \text{falls } (\mathcal{A}(t_1), \dots, \mathcal{A}(t_k)) \in P^{\mathcal{A}} \\ 0, & \text{sonst} \end{cases}$$

- Falls F die Form $F = \neg G$ hat, so gilt:

$$\mathcal{A}(F) = \begin{cases} 1, & \text{falls } \mathcal{A}(G) = 0 \\ 0, & \text{sonst} \end{cases}$$

- Falls F die Form $F = (G \wedge H)$ hat, so gilt:

$$\mathcal{A}(F) = \begin{cases} 1, & \text{falls } \mathcal{A}(G) = 1 \text{ und } \mathcal{A}(H) = 1 \\ 0, & \text{sonst} \end{cases}$$

- Falls F die Form $F = (G \vee H)$ hat, so gilt:

$$\mathcal{A}(F) = \begin{cases} 1, & \text{falls } \mathcal{A}(G) = 1 \text{ oder } \mathcal{A}(H) = 1 \\ 0, & \text{sonst} \end{cases}$$

- Falls F die Form $F = \forall x \, G$ hat, so gilt:

$$\mathcal{A}(F) = \begin{cases} 1, & \text{falls f\"ur alle } d \in U_{\mathcal{A}} \text{ gilt: } \mathcal{A}_{[x/d]}(G) = 1 \\ 0, & \text{sonst} \end{cases}$$

- Falls F die Form $F = \exists x \, G$ hat, so gilt:

$$\mathcal{A}(F) = \begin{cases} 1, & \text{falls es ein } d \in U_{\mathcal{A}} \text{ gibt mit: } \mathcal{A}_{[x/d]}(G) = 1 \\ 0, & \text{sonst} \end{cases}$$

$\mathcal{A}_{[x/d]}$ bezeichne die Struktur $\mathcal{A}'$, die mit $\mathcal{A}$ identisch ist, bis auf die Definition von $x^{\mathcal{A}'}$: Es sei $x^{\mathcal{A}'} = d$, wobei $d \in U_{\mathcal{A}} = U_{\mathcal{A}'}$.

Man k\"onnte auch sagen: x wird mit dem Wert d belegt.

Analog zur Aussagenlogik definieren wir:

- Falls $\mathcal{A}(F) = 1$ für eine Formel F und einer zu F passenden Struktur $\mathcal{A}$ gilt, so schreiben wir $\mathcal{A} \models F$ und sagen: F *gilt* in $\mathcal{A}$ oder auch $\mathcal{A}$ ist *Modell* für F .
- Falls jede zu F passende Struktur ein Modell für F ist, so schreiben wir $\models F$ und sagen F ist (allgemein-) *gültig*.
- Falls es ein Modell für die Formel F gibt, so heißt F *erfüllbar*, andernfalls *unerfüllbar*.
- Weitere Begriffe wie *Folgerung* und *Äquivalenz* können aus der Aussagenlogik übertragen werden.

Analog zum aussagenlogischen Fall lässt sich zeigen, dass für jede Formel F gilt:

F ist genau dann (allgemein-)gültig, wenn $\neg F$ unerfüllbar ist.

Wir schieben für einen Augenblick die präzise mathematische Definition beiseite, damit vielleicht klarer wird, was wir eigentlich gerade tun.

Ersetzt man in einer Aussage A eine Konstante durch eine Variable x , so entsteht die *Aussageform* $A(x)$, auch *Prädikat* $A(x)$ genannt.

- Krefeld liegt am Rhein $\rightarrow x$ liegt am Rhein
- 5 ist eine natürliche Zahl $\rightarrow y$ ist eine natürliche Zahl

Häufig hat man weitere Variablen und zusätzlich auch Terme, die über Variablen gebildet werden, z.B.

- 7 ist kleiner als 12 $\rightarrow x$ ist kleiner als y
- $\rightarrow$ Wir erhalten ein Prädikat der Form $A(x, y)$.
- 7 plus 3 ist kleiner als 12 $\rightarrow x + y$ ist kleiner als z
- $\rightarrow$ Wir erhalten ein Prädikat der Form $A(t(x, y), z)$.

- Eine Aussageform/ein Prädikat hat im Allgemeinen keinen bestimmten Wahrheitswert.
- Erst wenn die Variable durch einen konkreten Wert ersetzt wird, entsteht eine Aussage, die einen Wahrheitswert hat.

Aus der Aussagenlogik wird durch die Aussageformen die *Prädikatenlogik*. Im Folgenden einige Beispiele dazu.

Aussageform: $A(x) := x^2 > 30$

- $A(x)$ wird für $x = 6$ zur wahren Aussage.
- Für $x = 5$ ist $A(x)$ falsch.
- $\neg A(x)$ lautet $x^2 \leq 30$.

Aussageform: $A(x) := x$ ist eine Ampelfarbe

- Für $x \in \{\text{rot, gelb, grün}\}$ wird $A(x)$ wahr.
- Für $x = \text{blau}$ wird $A(x)$ zu einer falschen Aussage.

Der *Allquantor* $\forall$ steht für den Text „für alle“.

- $\forall x \in E : A(x)$ ist die Aussage, die in Textform lautet:
„Für alle Elemente x von E gilt: Die Aussageform $A(x)$ wird eine wahre Aussage.“
Oder anders formuliert: $A(x)$ wird für jedes $x \in E$ wahr.

Beispiele für wahre Aussagen:

- $\forall x \in \{\text{rot, gelb, grün}\} : x$ ist eine Ampelfarbe
- $\forall x \in \{-3, -2, -1, 0, 1, 2, 3\} : (x^2 = 4 \iff x = 2 \vee x = -2)$

Der *Existenzquantor* $\exists$ steht für den Text „es existiert“.

- $\exists x \in E : A(x)$ ist die Aussage, die in Textform lautet:
„Es existiert (mindestens) ein Element x von E , so dass $A(x)$ für dieses Element x wahr wird.“
Oder anders formuliert: $A(x)$ wird für (mindestens) ein $x \in E$ wahr.

Beispiele für wahre Aussagen:

- $\exists x \in \{\text{blau, gelb, grün}\} : x \text{ ist eine Ampelfarbe}$
- $\exists x \in \{1, 2, 3\} : x^2 = 4$

Das Prädikat $mutter(a, b)$ bedeute, dass a Mutter von b ist. Beschreiben Sie umgangssprachlich die folgenden Formeln:

- $\forall y \exists x \text{ mutter}(x, y)$
- $\exists x \forall y \text{ mutter}(x, y)$

Drücken Sie folgende Sachverhalte mittels Prädikaten aus:

- Alle Studenten sind arm.
- Alle grünen Pilze sind giftig.
- Es gibt einen armen Studenten.
- Wer schneller als Anna ist, ist auch schneller als Klaus.
- Niemand ist schneller als Harald.
- Claudia ist die schnellste.

Formulieren Sie das Kannibalen und Missionare-Problem mittels Prädikatenlogik:

- 3 Missionare und 3 Kannibalen wollen einen Fluss überqueren.
- Es steht nur ein Ruderboot für die Überfahrt zur Verfügung.
- Das Boot kann maximal 2 Personen aufnehmen.
- Es dürfen nie mehr Kannibalen als Missionare an einem Ort sein, sonst gibt es ein großes Fressen.

Prädikatenlogik

- Grundbegriffe
- *Normalformen*
- Unentscheidbarkeit
- Herbrand-Theorie
- Resolution

Zwei prädikatenlogische Formeln F und G heißen *äquivalent*, falls für alle sowohl zu F als auch zu G passenden Strukturen $\mathcal{A}$ gilt:
 $\mathcal{A}(F) = \mathcal{A}(G)$

Alle in der Aussagenlogik gezeigten Äquivalenzen gelten auch in der Prädikatenlogik:

- Idempotenz
- Kommutativität
- Assoziativität
- Absorption
- Distributivität
- Doppelnegation
- deMorgansche Regeln
- usw.

Um auch prädikatenlogische Formeln in Normalformen umformen zu können, benötigen wir noch Äquivalenzen, die Quantoren mit einbeziehen.

Seien F und G beliebige Formeln. Dann gilt:

$$1.a \quad \neg \forall x F \equiv \exists x \neg F$$

Beispiel: „Nicht alle Fußballspiele der 1. Bundesliga am vergangenen Spieltag endeten unentschieden.“

Ist äquivalent zu: „Es gab am vergangenen Spieltag der 1. Bundesliga (mindestens) ein Fußballspiel, das nicht unentschieden endete.“

Beispiel: „Nicht alle Außerirdischen haben grüne Haut.“

Ist äquivalent zu: „Es gibt Außerirdische, die keine grüne Haut haben.“

$$1.b \quad \neg \exists x F \equiv \forall x \neg F$$

Beispiel: „Es gab kein Fußballspiel der 1. Bundesliga, das am vergangenen Spieltag unentschieden endete.“

Ist äquivalent zu: „Alle Fußballspiele der 1. Bundesliga des vergangenen Spieltags endeten nicht unentschieden.“

2. Falls x in G nicht frei vorkommt, gilt:

$$(\forall x \, F \wedge G) \equiv \forall x \, (F \wedge G)$$

$$(\forall x \, F \vee G) \equiv \forall x \, (F \vee G)$$

$$(\exists x \, F \wedge G) \equiv \exists x \, (F \wedge G)$$

$$(\exists x \, F \vee G) \equiv \exists x \, (F \vee G)$$

3. $(\forall x \, F \wedge \forall x \, G) \equiv \forall x \, (F \wedge G)$

$$(\exists x \, F \vee \exists x \, G) \equiv \exists x \, (F \vee G)$$

4. $\forall x \forall y \, F \equiv \forall y \forall x \, F$

$$\exists x \exists y \, F \equiv \exists y \exists x \, F$$

Übung: Formulieren Sie Beispiele für obige Aussagen.

Exemplarisch beweisen wir: $(\forall x F \wedge G) \equiv \forall x (F \wedge G)$

$$\mathcal{A}(\forall x F \wedge G) = 1$$

$$\iff \mathcal{A}(\forall x F) = 1 \text{ und } \mathcal{A}(G) = 1$$

$$\iff \text{Für alle } d \in U_{\mathcal{A}} \text{ gilt } \mathcal{A}_{[x/d]}(F) = 1 \text{ und } \mathcal{A}_{[x/d]}(G) = 1, \text{ denn} \\ \text{weil } x \text{ in } G \text{ nicht frei vorkommt, gilt } \mathcal{A}(G) = \mathcal{A}_{[x/d]}(G).$$

$$\iff \text{Für alle } d \in U_{\mathcal{A}} \text{ gilt: } \mathcal{A}_{[x/d]}((F \wedge G)) = 1$$

$$\iff \mathcal{A}(\forall x (F \wedge G)) = 1$$

Übung: Zeigen Sie, dass folgende Formelpaare *nicht* äquivalent sind, indem Sie Gegenbeispiele angeben:

$$(\forall x F \vee \forall x G) \not\equiv \forall x (F \vee G)$$

$$(\exists x F \wedge \exists x G) \not\equiv \exists x (F \wedge G)$$

Übung: Bei verschiedenen Quantoren ist die Reihenfolge extrem wichtig. Es gilt folgende Nicht-Äquivalenz:

$$\forall x \exists y F \not\equiv \exists y \forall x F$$

Formulieren Sie ein Beispiel, das den Sachverhalt deutlich macht.

Übung: Zeigen Sie, dass $\forall x \exists y P(x, y)$ eine Folgerung von $\exists y \forall x P(x, y)$ ist, aber nicht umgekehrt.

Durch obige Äquivalenzumformungen können wir die in einer Formel evtl. vorhandenen Quantoren an den Anfang der Formel verschieben:

$$\begin{aligned} & (\neg(\exists x P(x, y) \vee \forall z Q(z)) \wedge \exists w P(f(a, w))) \\ \equiv & ((\neg\exists x P(x, y) \wedge \neg\forall z Q(z)) \wedge \exists w P(f(a, w))) \text{ deMorgan} \\ \equiv & ((\forall x \neg P(x, y) \wedge \exists z \neg Q(z)) \wedge \exists w P(f(a, w))) \text{ wegen 1.} \\ \equiv & (\exists w P(f(a, w)) \wedge (\forall x \neg P(x, y) \wedge \exists z \neg Q(z))) \text{ kommutativ} \\ \equiv & \exists w (P(f(a, w)) \wedge \forall x (\neg P(x, y) \wedge \exists z \neg Q(z))) \text{ wegen 2.} \\ \equiv & \exists w (\forall x (\exists z \neg Q(z) \wedge \neg P(x, y)) \wedge P(f(a, w))) \text{ kommutativ} \\ \equiv & \exists w (\forall x \exists z (\neg Q(z) \wedge \neg P(x, y)) \wedge P(f(a, w))) \text{ wegen 2.} \\ \equiv & \exists w \forall x \exists z ((\neg Q(z) \wedge \neg P(x, y)) \wedge P(f(a, w))) \text{ wegen 2.} \end{aligned}$$

Punkt 2 des obigen Satzes kann nicht immer angewendet werden:

$$\forall x P(x, y) \wedge Q(x) \not\equiv \forall x (P(x, y) \wedge Q(x))$$

Um die Regel 2 dennoch anwenden zu können, benennen wir zunächst die freie Variable x um:

$$\begin{aligned}\forall x P(x, y) \wedge Q(x) &\equiv \forall x P(x, y) \wedge Q(z) \\ &\equiv \forall x (P(x, y) \wedge Q(z))\end{aligned}$$

Ein solches Umbenennen von Variablen heißt *Substitution*. Sei F eine Formel, x eine Variable und t ein Term.

Dann bezeichnet $F[x/t]$ diejenige Formel, die man aus F erhält, indem jedes freie Vorkommen der Variablen x in F durch den Term t ersetzt wird.

Anmerkung: Wir müssen zwischen $\mathcal{A}_{[x/d]}$ und $F[x/t]$ deutlich unterscheiden. Im ersten Fall erhalten wir eine neue Struktur, bei der x den Wert d erhält (Semantik), im zweiten Fall erhalten wir eine neue Formel, bei der x durch t ersetzt ist (Syntax).

Lemma: (gebundene Umbenennung)

Falls die Variable y nicht in der Formel G vorkommt, gilt:

$$\forall x \ G \equiv \forall y \ G[x/y]$$

$$\exists x \ G \equiv \exists y \ G[x/y]$$

Beispiel: Aus

$$\forall x \ P(x) \wedge Q(x, y)$$

wird durch gebundene Umbenennung die äquivalente Formel

$$\forall z \ P(z) \wedge Q(x, y).$$

Durch systematisches Ausführen von gebundenen Umbenennungen, wobei immer neue, noch nicht vorkommende Variablen verwendet werden, kann man das folgende Lemma zeigen.

Lemma: Zu jeder Formel F gibt es eine äquivalente Formel G in bereinigter Form.

Eine Formel heißt *bereinigt*, sofern es keine Variable gibt, die in der Formel sowohl frei als auch gebunden vorkommt, und sofern hinter allen vorkommenden Quantoren verschiedene Variablen stehen.

Beispiel:

$$\begin{aligned} & \forall x \exists y P(x, f(y)) \wedge \forall y (Q(x, y) \vee R(x)) \\ & \equiv \forall u \exists y P(u, f(y)) \wedge \forall y (Q(x, y) \vee R(x)) \\ & \equiv \forall u \exists v P(u, f(v)) \wedge \forall y (Q(x, y) \vee R(x)) \end{aligned}$$

Freie Variablen: x

Gebundene Variablen: u, v, y

Eine Formel heißt *pränex* oder in *Pränexform*, falls sie die Bauart

$$Q_1 y_1 \ Q_2 y_2 \ \dots \ Q_n y_n \ F$$

hat, wobei $Q_i \in \{\exists, \forall\}$, $n \geq 0$, und die y_i Variablen sind. Ferner kommt kein Quantor in F vor.

Satz: Für jede Formel F gibt es eine äquivalente und bereinigte Formel G in Pränexform. (Beweis: Induktion über Formelaufbau.)

Beispiel:

$$\begin{aligned} & \forall x \ \exists y \ [P(x, g(y, f(x))) \ \vee \ \neg Q(z)] \ \vee \ \neg \forall x \ R(x, y) \\ & \equiv \forall x \ \exists y \ [P(x, g(y, f(x))) \ \vee \ \neg Q(z)] \ \vee \ \exists x \neg R(x, y) \\ & \equiv \forall x \ \exists y \ [P(x, g(y, f(x))) \ \vee \ \neg Q(z)] \ \vee \ \exists v \neg R(v, y) \\ & \equiv \forall x \ \exists u \ [P(x, g(u, f(x))) \ \vee \ \neg Q(z)] \ \vee \ \exists v \neg R(v, y) \\ & \equiv \forall x \ \exists u \ \exists v \ [P(x, g(u, f(x))) \ \vee \ \neg Q(z) \ \vee \ \neg R(v, y)] \end{aligned}$$

Für die prädikatenlogische Variante der Resolution müssen wir die Formeln noch weiter vereinfachen: Existenzquantoren werden durch *Skolemisierung* eliminiert.

- Tritt die Form $\exists x$ nicht im Geltungsbereich eines Allquantors auf, wird $\exists x$ einfach weggelassen und alle Auftreten von x werden durch eine neue Konstante c ersetzt.

Beispiel: Aus $\exists x P(x)$ wird $P(c)$.

Beispiel: Aus $\exists x \exists z \forall y P(y, x, z)$ wird $\forall y P(y, c_1, c_2)$.

- Tritt dagegen $\exists x$ im Geltungsbereich der allquantifizierten Variablen $y_1, \dots, y_k$ auf, so werden alle Auftreten von x durch den Term $f(y_1, \dots, y_k)$ ersetzt, wobei f ein neues Funktionssymbol sein muss.

Beispiel: Aus $\forall y \exists x P(x, y)$ wird $\forall y P(f(y), y)$.

Beispiel: Aus $\exists x \forall y \exists z P(y, x, z)$ wird $\forall y P(y, c, f(y))$.

Die vollständige Skolemisierung einer Formel F entfernt alle Existenzquantoren aus F und liefert eine Formel F' , die erfüllbarkeitsäquivalent zu F ist, d.h. F ist genau dann erfüllbar, wenn F' erfüllbar ist.

Übung: Geben Sie die Skolemform der folgenden Formel an:

$$F = \forall x \exists y \forall z \exists w (\neg P(a, w) \vee Q(f(x), y))$$

Übung: Wenden Sie alle vorgestellten Umformungsschritte (bereinigen, Pränexform, Skolemform) an auf die Formel:

$$F = \forall z \exists y (P(x, g(y), z) \vee \neg \forall x Q(x)) \wedge \neg \forall z \exists x \neg R(f(x, z), z)$$

Gegeben sei eine prädikatenlogische Formel F .

- Berechne F durch systematisches Umbenennen der gebundenen Variablen.
→ Es entsteht eine zu F äquivalente Formel F_1 .
- Binde die in F_1 vorkommenden freien Variablen $y_1, \dots, y_k$ (ersetze F_1 durch $F_2 = \exists y_1 \exists y_2 \dots \exists y_k F_1$)
→ Dann ist F_2 erfüllbarkeitsäquivalent zu F_1 .
- Stelle eine zu F_2 äquivalente Formel F_3 in Pränexform her.
→ F_3 ist erfüllbarkeitsäquivalent zu F .
- Eliminiere die in F_3 vorkommenden Existenzquantoren durch Skolemisierung zur Formel F_4 .
→ F_4 ist erfüllbarkeitsäquivalent zu F_3 und damit auch zu F .

- Da alle auftretenden Variablen jetzt allquantifiziert sind, können (zur Darstellung) alle Allquantoren weggelassen werden.
- Die so erhaltene Formel kann durch entsprechende Umformungen in die disjunktive oder die konjunktive Normalform überführt werden und in *Klauselform* dargestellt werden.

Prädikatenlogik

- Grundbegriffe
- Normalformen
- *Unentscheidbarkeit*
- Herbrand-Theorie
- Resolution

Wir suchen einen algorithmischen Test für die Erfüllbarkeit oder Gültigkeit von prädikatenlogischen Formeln.

Gültigkeitsproblem der Prädikatenlogik:

gegeben: Eine prädikatenlogische Formel F .

gefragt: Ist F eine (allgemein-)gültige Formel?

Erfüllbarkeitsproblem der Prädikatenlogik:

gegeben: Eine prädikatenlogische Formel F .

gefragt: Ist F erfüllbar?

Leider können solche Algorithmen aus prinzipiellen Gründen nicht existieren, wie wir im Folgenden sehen werden.

Postisches Korrespondenzproblem:

gegeben: Eine endliche Menge von Wortpaaren $(x_1, y_1), (x_2, y_2), \dots, (x_k, y_k)$, wobei $x_i, y_i \in \Sigma^+$.

gefragt: Gibt es eine Folge von Indizes $i_1, i_2, \dots, i_n \in \{1, 2, \dots, k\}$, $n \geq 1$, mit $x_{i_1} x_{i_2} \dots x_{i_n} = y_{i_1} y_{i_2} \dots y_{i_n}$?

Beispiel: $K = ((1, 101), (10, 00), (011, 11))$ besitzt die Lösung $(1, 3, 2, 3)$, denn es gilt

$$\underbrace{1}_{x_1} \underbrace{011}_{x_3} \underbrace{10}_{x_2} \underbrace{011}_{x_3} = \underbrace{101}_{y_1} \underbrace{11}_{y_3} \underbrace{00}_{y_2} \underbrace{11}_{y_3}$$

Aus der Berechenbarkeitstheorie wissen wir, dass das Postische Korrespondenzproblem nicht entscheidbar ist.

Mittels Reduktion kann man zeigen, dass das Gültigkeitsproblem für prädikatenlogische Formeln nicht entscheidbar ist.

Als Folgerung ergibt sich, dass auch das Erfüllbarkeitsproblem für prädikatenlogische Formeln nicht entscheidbar ist. (Eine Formel F ist gültig, genau dann wenn $\neg F$ unerfüllbar ist.)

Monadische Prädikatenlogik: Die Formeln enthalten keine Funktionssymbole und alle Prädikatsymbole sind einstellig.

- Für eine erfüllbare Aussage F der monadischen Prädikatenlogik mit den einstelligen Prädikatensymbolen $P_1, \dots, P_n$ existiert bereits ein Modell der Mächtigkeit 2^n .
- Das Erfüllbarkeits- und das Gültigkeitsproblem für monadische prädikatenlogische Formeln ist entscheidbar.

Prädikatenlogik

- Grundbegriffe
- Normalformen
- Unentscheidbarkeit
- *Herbrand-Theorie*
- Resolution

Problem im Umgang mit prädikatenlogischen Formeln:

- Die Definition von Strukturen $\mathcal{A} = (U_{\mathcal{A}}, I_{\mathcal{A}})$ lässt es zu, dass $U_{\mathcal{A}}$ beliebige Mengen sein können.
 - Es scheint keine Möglichkeit zu geben, Aussagen über deren Mächtigkeit, geschweige denn über den Aufbau von $I_{\mathcal{A}}$ zu machen.
- Wie sollen wir alle potenziellen Strukturen für eine gegebene Formel auf systematische Art und Weise aufzählen, um sie daraufhin auf ihre Modelleigenschaft zu untersuchen?

Arbeiten, die im wesentlichen auf Jacques Herbrand, Kurt Gödel und Thoralf Skolem zurückgehen, zeigen: Die algorithmische Suche nach potenziellen Modellen für eine Formel kann auf eine gewisse kanonische Art erfolgen bzw. beschränkt bleiben.

Sei F eine geschlossene Formel in Skolemform. Das *Herbrand-Universum* $D(F)$ wird wie folgt induktiv definiert:

- $D(F)$ enthält alle in F vorkommenden Konstanten. Falls F keine Konstanten enthält, so ist a in $D(F)$.
- Für jedes in F vorkommende n -stellige Funktionssymbol f und Terme $t_1, t_2, \dots, t_n$ in $D(F)$ ist auch der Term $f(t_1, t_2, \dots, t_n)$ in $D(F)$.

Mit anderen Worten: Das Herbrand-Universum besteht aus allen variablenfreien Termen, die aus den Konstanten und den Funktionssymbolen der Formel F gebildet werden können.

Beispiel: Das Herbrand-Universum zur Formel $F = \forall x \forall y \forall z P(x, f(y), g(z, x))$ ist:

$$D(F) = \{a, f(a), g(a, a), f(f(a)), f(g(a, a)), \\ g(f(a), a), g(a, f(a)), g(f(a), f(a)), \dots\}$$

Übung: Geben Sie für die Formel $G = \forall x \forall y Q(c, f(x), h(y, b))$ die ersten Elemente des Herbrand-Universums an.

Wir werden im Folgenden sehen, dass es ausreichend ist, $D(F)$ als Grundbereich zu nutzen, um nach möglichen Modellen für die Formel F zu suchen.

Sei F eine Aussage in Skolemform. Dann heißt jede zu F passende Struktur $\mathcal{A} = (U_{\mathcal{A}}, I_{\mathcal{A}})$ eine *Herbrand-Struktur* für F , falls gilt:

- $U_{\mathcal{A}} = D(F)$
- Für jedes in F vorkommende n -stellige Funktionssymbol f und $t_1, t_2, \dots, t_n \in D(F)$ ist $f^{\mathcal{A}}(t_1, t_2, \dots, t_n) = f(t_1, t_2, \dots, t_n)$.

Beispiel: Eine Herbrand-Struktur $\mathcal{A} = (U_{\mathcal{A}}, I_{\mathcal{A}})$ für die Formel $F = \forall x \forall y \forall z P(x, f(y), g(z, x))$ muss folgende Bedingungen erfüllen:

$$U_{\mathcal{A}} = D(F) = \{a, f(a), g(a, a), f(f(a)), f(g(a, a)), \dots\}$$

und

$$\begin{aligned} f^{\mathcal{A}}(a) &= f(a) \\ f^{\mathcal{A}}(f(a)) &= f(f(a)) \\ f^{\mathcal{A}}(g(a, a)) &= f(g(a, a)) \end{aligned}$$

usw.

Die Wahl von $P^{\mathcal{A}}$ ist noch offen. Wir könnten festlegen:

$$(t_1, t_2, t_3) \in P^{\mathcal{A}} \iff g(t_1, t_2) = g(t_2, f(t_3))$$

Ist die so definierte Struktur $\mathcal{A}$ ein Modell für F ?

Anmerkungen:

- Bei Herbrand-Strukturen sind der Grundbereich und die Interpretation der Funktionssymbole festgelegt. Nur die Interpretation der Prädikatsymbole kann frei gewählt werden.
- Syntax und Semantik von Termen wird „gleichgeschaltet“. Terme werden „durch sich selbst“ interpretiert.
- Eine Herbrand-Struktur $\mathcal{A}$ für eine Formel F nennen wir *Herbrand-Modell* für F , falls $\mathcal{A}$ ein Modell für F ist.

Satz: Sei F eine Aussage in Skolemform. Dann ist F genau dann erfüllbar, wenn F ein Herbrand-Modell besitzt. (ohne Beweis)

Folgerung: Jede erfüllbare Formel der Prädikatenlogik besitzt bereits ein abzählbares Modell, also ein Modell mit abzählbarer Grundmenge.

Sei $F = \forall y_1 \forall y_2 \dots \forall y_n F^*$ eine Aussage in Skolemform. Dann ist $E(F)$, die *Herbrand-Expansion* von F , definiert als

$$E(F) = \{F^*[y_1/t_1][y_2/t_2] \cdots [y_n/t_n] \mid t_1, t_2, \dots, t_n \in D(F)\}.$$

Die Formeln in $E(F)$ entstehen also, indem die Terme in $D(F)$ in jeder möglichen Weise für die Variablen in F^* substituiert werden.

Beispiel: Für die obige Formel $F = \forall x \forall y \forall z P(x, f(y), g(z, x))$ sind die einfachsten Formeln in $E(F)$ die folgenden:

$$\begin{array}{ll} P(a, f(a), g(a, a)) & \text{mit } [x/a][y/a][z/a] \\ P(f(a), f(a), g(a, f(a))) & \text{mit } [x/f(a)][y/a][z/a] \\ P(a, f(f(a)), g(a, a)) & \text{mit } [x/a][y/f(a)][z/a] \\ P(a, f(a), g(f(a), a)) & \text{mit } [x/a][y/a][z/f(a)] \\ P(g(a, a), f(a), g(a, g(a, a))) & \text{mit } [x/g(a, a)][y/a][z/a] \end{array}$$

USW.

Beachte: Die Formeln in $E(F)$ können letztlich als aussagenlogische Formeln behandelt werden, da sie keine Variablen enthalten.

Satz: Für jede Aussage F in Skolemform gilt: F ist erfüllbar genau dann, wenn die Formelmenge $E(F)$ im aussagenlogischen Sinn erfüllbar ist. (ohne Beweis)

Zusammenfassend kann man sagen:

- Eine prädikatenlogische Formel kann durch i.Allg. unendlich viele aussagenlogische Formeln approximiert werden.
- Diese aussagenlogischen Formeln können systematisch aufgezählt werden.

Um daraus ein Semi-Entscheidungsverfahren für die Unerfüllbarkeit von Formeln der Prädikatenlogik zu formulieren, benötigen wir noch den Endlichkeitssatz der Aussagenlogik.

Endlichkeitssatz: (Aussagenlogik)

Eine Menge $\mathcal{F}$ von aussagenlogischen Formeln ist genau dann erfüllbar, wenn jede endliche Teilmenge von $\mathcal{F}$ (also auch $\mathcal{F}$ selbst) erfüllbar ist. (ohne Beweis)

Oder anders formuliert: Eine evtl. unendliche Formelmenge $\mathcal{F}$ ist genau dann unerfüllbar, wenn bereits eine endliche Teilmenge $\mathcal{F}'$ von $\mathcal{F}$ unerfüllbar ist.

Satz von Herbrand:

Eine Aussage F in Skolemform ist genau dann unerfüllbar, wenn es eine endliche Teilmenge von $E(F)$ gibt, die im aussagenlogischen Sinn unerfüllbar ist. (direkte Folgerung aus obigen Sätzen)

Algorithmus von Gilmore: Semi-Entscheidungsverfahren für die Unerfüllbarkeit von Formeln der Prädikatenlogik.

- Gegeben: Eine prädikatenlogische Aussage F in Skolemform.
- Sei $E(F) = \{F_1, F_2, F_3, \dots\}$ eine beliebige aber fixierte Aufzählung der Herbrand-Expansion von F .
- Algorithmus:
 - $n := 0$;
 - repeat $n := n + 1$;
 - until $(F_1 \wedge F_2 \wedge \dots \wedge F_n)$ ist unerfüllbar

Anmerkungen:

- semi-entscheidbar $\hat{=}$ rekursiv aufzählbar
- Für unerfüllbare Formeln stoppt obiger Algorithmus, für erfüllbare Formeln dagegen nicht.
- Durch Eingabe von $\neg F$ anstelle von F wird aus dem Unerfüllbarkeitstest ein Gültigkeitstest.

Übung: Zeigen Sie mittels des Algorithmus von Gilmore, dass folgende Formeln unerfüllbar sind:

- $F = \exists x \forall y (\neg P(f(f(x))) \wedge P(f(y)))$
- $G = \forall x (P(f(x)) \wedge \neg P(x))$

Anmerkung: Der Algorithmus von Gilmore funktioniert zwar, ist aber in der Praxis unbrauchbar, weil er zu viele Formeln erzeugt und nicht zielgerichtet arbeitet.

Daher werden wir uns im Folgenden ansehen, wie wir das Prinzip der Resolution auch in der Prädikatenlogik anwenden können.

Prädikatenlogik

- Grundbegriffe
- Normalformen
- Unentscheidbarkeit
- Herbrand-Theorie
- *Resolution*

Die Tests auf Unerfüllbarkeit können auch per aussagenlogischer Resolution erledigt werden.

Beispiel: Gegeben sei folgende unerfüllbare Formel in Skolemform:

$$F = \forall x (P(x) \wedge \neg P(f(x)))$$

Damit erhalten wir

- als Klauselform: $\{\{P(x)\}, \{\neg P(f(x))\}\}$
- als Herbrand-Universum: $D(F) = \{a, f(a), f(f(a)), \dots\}$
- als Herbrand-Expansion:
 $E(F) = \{ (P(a) \wedge \neg P(f(a))), (P(f(a)) \wedge \neg P(f(f(a)))) , \dots \},$
oder in Klauselform
 $\{\{P(a)\}, \{\neg P(f(a))\}, \{P(f(a))\}, \{\neg P(f(f(a)))\}, \dots\}$

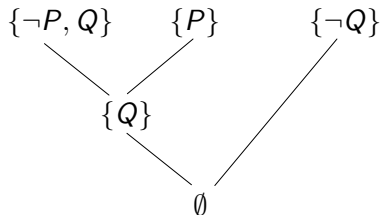
Bereits die ersten beiden Elemente von $E(F)$ reichen aus, um die Unerfüllbarkeit von F zu zeigen:

$$\begin{array}{ccccccc} \{P(a)\} & & \{\neg P(f(a))\} & & \{P(f(a))\} & & \{\neg P(f(f(a)))\} \\ & & \searrow & & \swarrow & & \\ & & \emptyset & & & & \end{array}$$

Um besser zu verstehen, was wir hier eigentlich tun, werfen wir kurz einen Blick zurück auf die Aussagenlogik.

Aussagenlogik:

- Betrachten wir die Klauselmenge $\{\{\neg P, Q\}, \{P\}, \{\neg Q\}\}$.
- Resolutionsprinzip: Resolviere Klauseln mit komplementären Literalen.



- Führt der Resolutionsgraph auf die leere Klausel, dann ist die Formel unerfüllbar.

Aber was passiert, wenn wir aus den Variablen P und Q der obigen Formel Prädikate machen?

- Betrachten wir bspw. folgende Formel:

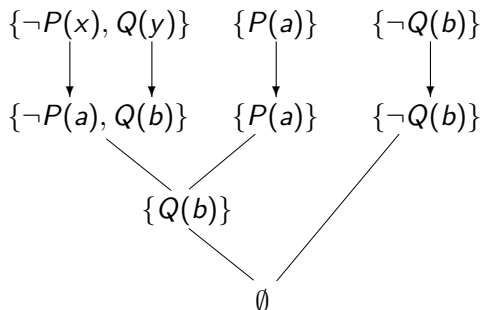
$$\forall x \forall y (P(x) \rightarrow Q(y)) \wedge P(a) \wedge \neg Q(b)$$

- Als Klauselmenge: $\{\{\neg P(x), Q(y)\}, \{P(a)\}, \{\neg Q(b)\}\}$

Sind $\neg P(x)$ und $P(a)$ komplementär?

- In einem gewissen Sinn: Ja!
- Denn obige Formel soll für alle x gelten, also insbesondere auch für $x := a$.

Im Prinzip suchen wir also Substitutionen (hier $[x/a]$ und $[y/b]$), die komplementäre „Litere“ erzeugen.



Die Herbrand-Theorie zeigt formal die Korrektheit des obigen Vorgehens und erweitert das obige Vorgehen auch auf Funktionen bzw. Terme.

Übung: Zeigen Sie mittels Resolution, dass $\exists y \text{ irren}(y)$ eine Folgerung ist von:

$$(\forall x \text{ Mensch}(x) \rightarrow \text{irren}(x)) \wedge \text{Mensch}(\text{Max})$$

Oder in Worten: Alle Menschen irren sich und Max ist ein Mensch, also gibt es jemanden, der sich irrt.

Problem:

Wie sollen wir effizient die Substitutionen finden, die bei einer Resolutionsherleitung zur leeren Klausel führen? Ein systematisches Durchprobieren aller Substitutionen scheint sehr aufwendig und wenig zielgerichtet zu sein.

Problem: Es müssen Entscheidungen für einige Substitutionen „vorausschauend“ getroffen werden, um Resolutionen, die im Diagramm weiter unten liegen, zu ermöglichen.

Wir wollen daher nur insoweit Substitutionen durchführen, wie es für den direkt nachfolgenden Resolutionsschritt notwendig ist.

Beispiel: Bei $\{P(x), \neg Q(g(x))\}$ und $\{\neg P(f(y))\}$ reicht die Substitution $[x/f(y)]$, um die Resolvente $\{\neg Q(g(f(y)))\}$ zu erhalten. Es ist nicht nötig, hier bereits y zu substituieren.

Wir suchen also Substitutionen, die zwei Literale (oder mehr) *unifiziert*, also identisch macht. Im obigen Beispiel unifiziert $[x/f(y)]$ die beiden Literale $P(x)$ und $P(f(y))$.

Auch die Substitution $[x/f(a)][y/a]$ würde die beiden Literale unifizieren, aber dabei wird mehr substituiert als notwendig.

Eine Substitution σ ist ein *Unifikator* einer (endlichen) Menge von Literalen $\mathcal{L} = \{L_1, L_2, \dots, L_k\}$, falls $L_1\sigma = L_2\sigma = \dots = L_k\sigma$. Wir sagen dann: $\mathcal{L}$ ist *unifizierbar*.

Ein Unifikator σ einer Literalmenge $\mathcal{L}$ heißt *allgemeinster Unifikator* von $\mathcal{L}$, falls für jeden Unifikator μ von $\mathcal{L}$ gilt, dass es eine Substitution s gibt mit $\mu = \sigma s$.

Beispiel: Betrachten wir die Terme $f(x, g(a, y))$ und $f(b, z)$. Dann sind σ und μ zwei Unifikatoren.

$$\mu = [x/b][y/a][z/g(a, a)] \qquad \sigma = [x/b][z/g(a, y)]$$

σ ist allgemeinster Unifikator, μ setzt sich dagegen aus σ und einer anderen Substitution zusammen:

$$\mu = \underbrace{[x/b][z/g(a, a)]}_{\sigma} [y/a]$$

Beispiele: Der allgemeinste Unifikator (most general unifier) von

- $\{P(x), P(f(y))\}$ ist $[x/f(y)]$.
- $\{Q(x, f(x)), Q(y, g(y))\}$ existiert nicht.
- $\{P(x), P(f(x))\}$ existiert nicht.
- $\{Q(x, f(y)), Q(g(z), f(x))\}$ ist $[x/g(z)][y/g(z)]$.

Anmerkungen:

- Der allgemeinste Unifikator zweier Terme ist eindeutig bestimmt (bis auf Variablenumbenennungen).
- Eine Menge $\mathcal{L}$ von (mehr als zwei) Literalen kann nicht unifizierbar sein, auch wenn alle Paare von Literalen in $\mathcal{L}$ unifizierbar sind.

Beispiel: $\{Q(x, f(y)), Q(f(y), z), Q(z, f(x))\}$

Unifikationsalgorithmus:

Eingabe: eine Literalmenge $\mathcal{L}$

$\sigma := []$; (leere Substitution)

while $|\mathcal{L}\sigma| > 1$ *do*

 Suche die erste Position in $\mathcal{L}\sigma$, an der sich zwei Literale
 L_1, L_2 unterscheiden

if keines der beiden Zeichen ist eine Variable

then stoppe mit „nicht unifizierbar“

else Sei x die Variable und t der Term in anderen Literal
 (möglicherweise auch eine Variable)

if x kommt in t vor

then stoppe mit „nicht unifizierbar“

else $\sigma := \sigma[x/t]$

Ausgabe: σ

Anmerkung: Die Anzahl der Resolutionsmöglichkeiten ist im Allgemeinen sehr hoch, weil

- es viele Möglichkeiten gibt, zwei Klauseln zu finden, die resolvierbar sind, um so neue Resolventen zu erzeugen.
- Resolventen bei einer Resolutionsherleitung länger werden können als ihre Ausgangsklauseln.

Für die Resolution der leeren Klausel eignen sich aber eventuell nur sehr wenige Paare von Klauseln. Wie finden wir also die „richtigen“ Klauseln zum Resolvieren?

Bei der logischen Programmierung beschränken wir uns auf Hornklauseln:

- Jede Programmklausel (also eine Regel oder ein Fakt) enthält genau ein positives Literal (definite Klausel).
- Eine Zielklausel (Anfrage an das System) enthält nur negative Literale.

Das hat weitreichende Folgen für die Ableitungsmöglichkeiten und die dabei erzielbare Effizienz:

- Resolvieren zweier Programmklauseln kann niemals zur leeren Klausel führen.
 - Die Resolution einer Zielklausel und einer Programmklausel führt immer wieder zu einer neuen (evtl. leeren) Zielklausel.
- SLD-Resolution: Selection rule driven Linear resolution for Definite clauses

Linear resolution bedeutet, dass mit der gerade hergeleiteten Klausel weitergearbeitet wird.

Die hier vorgestellte Prädikatenlogik nennt man *Prädikatenlogik erster Stufe* oder auch *Prädikatenlogik erster Ordnung*.

- PL1 erweitert die Aussagenlogik um Quantoren, Prädikate und Terme.
- Leider lassen sich mit PL1 viele wichtige Formeln nicht ausdrücken, bspw. das Induktionsaxiom der Peano-Axiome:

$$\forall P (P(0) \wedge (\forall x P(x) \rightarrow P(x + 1))) \rightarrow \forall x P(x)$$

- Die *Prädikatenlogik zweiter Ordnung* erlaubt Quantifizierungen auch über Funktions- und Prädikatsymbole.
- Das macht die Sache nicht leichter: Viele Eigenschaften, die in PL1 entscheidbar oder semi-entscheidbar sind, sind in PL2 unentscheidbar.

- Geschichte und Anwendungen
- Zustandsraumsuche
- Aussagen- und Prädikatenlogik
- *Einführung in Prolog*

- *Allgemeines*
- Erste Schritte
- Arithmetik
- Listen
- Backtracking
- Sortieren
- Input/Output
- Zahlenrätsel
- Kannibalen und Missionare
- 8-Damen-Problem
- Schiebepuzzle
- A^* -Algorithmus

Algorithm = Logic + Control

- C, C++ oder Java: Angabe eines Lösungsalgorithmus.
→ WIE ist ein Problem zu lösen?
 - Logikprogramme: Nur Angabe des Problems.
→ WAS ist zu lösen - nicht wie. (Idealfall)
 - Dieser Idealfall wird bei existierenden Prolog-Systemen mit depth-first Auswertungsstrategie nicht erreicht.
 - Die breadth-first Auswertungsstrategie ist zwar vollständig, aber hoffnungslos ineffizient.
- Wir müssen einen Kompromiss finden zwischen dem Idealfall (vollständige Trennung von Logik- und Kontrollkomponente) und Effizienz.

- Programmiersprache für Logik-Programme
- kennt zusätzlich Meta-Befehle, Ein-/Ausgabe und Arithmetik
- Im Gegensatz zu reinen Logik-Programmen wird das WIE durch die Verarbeitungsreihenfolge bestimmt.
- David H.D. Warren:
“Prolog ist [...] ein Werkzeug für das Denken.”
- Basiert auf Ideen von Kowalski und Colmerauer.
- Erste Implementierung 1972 in Marseille.
- Aufschwung durch 5th Generation-Projekt.
- Neben LISP die wichtigste Sprache der KI.
- Einsatzgebiete: Expertensysteme, Rule Engines
- Warren Abstract Machine: Virtuelle Machine, die kompilierten Prolog-Code ausführen kann (z.B. GNU Prolog).
- 1995 Standardisierung: ISO/IEC 13211-1, vorher viele Dialekte

Literatur

- W.F. Clocksin, C.S. Mellish:
Programming in Prolog.
Springer, New York, 1981
- L. Sterling, E. Shapiro:
Prolog - Fortgeschrittene Programmier Techniken.
Addison-Wesley, Bonn, 1988
- Ivan Bratko:
Prolog – Programming for Artificial Intelligence.
3rd edition, Addison-Wesley, 2000.

Es gibt syntaktische Konventionen, um verschiedene Programmbestandteile zu unterscheiden:

- Variablen: Großschreibung
- Prädikate und Funktionssymbole: Kleinschreibung
- Jede Klausel wird mit einem Punkt abgeschlossen.

In einer konkret einsetzbaren Sprache brauchen wir Konzepte, um Daten einzulesen, auszugeben oder zu bearbeiten: System stellt Prädikate wie `read` oder `write` bereit.

- Können vom Benutzer nicht verändert werden.
- Haben vom logischen Gesichtspunkt keine Bedeutung.
- Werden bei Bearbeitung durch Prologs Auswertemechanismus einfach übergangen.
- Erzeugen aber Seiteneffekte wie schreiben auf den Bildschirm oder in eine Datei.

Man muss sich über die Abfolge der Auswertung eines Prolog-Programms bewusst sein:

- Eine Zielklausel wie

`?- read(X), compute(X,Y), write(Y).`

kann sinnvoll nur von links nach rechts ausgewertet werden.

- Im Gegensatz zu den theoretischen Betrachtungen: Nur Klauseln ohne Seiteneffekte können auch in einer anderen Reihenfolge ausgewertet werden.

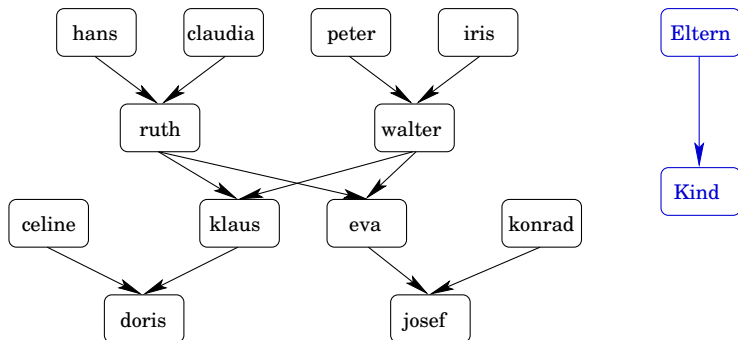
Andere Prädikate erzwingen Variableninstanziierungen:

- Trifft das Auswertesystem bspw. auf `is(X, Y*5)` (oder in Infix-Notation `X is Y*5`), wobei `Y` bereits mit der Konstanten 7 belegt ist, so wird `X` mit 35 belegt.

→ Auf diese Art können arithmetische Berechnungen in Prolog realisiert werden.

- Allgemeines
- *Erste Schritte*
- Arithmetik
- Listen
- Backtracking
- Sortieren
- Input/Output
- Zahlenrätsel
- Kannibalen und Missionare
- 8-Damen-Problem
- Schiebepuzzle
- A^* -Algorithmus

Wir wollen folgende Verwandschaftsbeziehung in Prolog darstellen:



Dazu definieren wir folgende Fakten:

<code>parent(hans, ruth).</code>	<code>parent(ruth, eva).</code>
<code>parent(claudia, ruth).</code>	<code>parent(walter, eva).</code>
<code>parent(peter, walter).</code>	<code>parent(celine, doris).</code>
<code>parent(iris, walter).</code>	<code>parent(klaus, doris).</code>
<code>parent(ruth, klaus).</code>	<code>parent(eva, josef).</code>
<code>parent(walter, klaus).</code>	<code>parent(konrad, josef).</code>

Dabei bezeichnet `parent(x, y)` den Sachverhalt, dass `y` das Kind von `x` ist. Wir lesen `parent(x, y)` also als `x is parent of y`.

Fakten sind Prädikate ohne Variablen.

Sie finden im Internet frei verfügbare Prolog-Interpreter z.B. unter:

- www.gprolog.org/
- jlogic.sourceforge.net/

Das Manual zu GNU-Prolog finden Sie unter

www.gprolog.org/manual/gprolog.html

Freie Prolog-Tutorials finden Sie bspw. unter:

- www.learnprolognow.org/
- www.csupomona.edu/~jrffisher/www/prolog_tutorial/
- ktiml.mff.cuni.cz/~bartak/prolog/

Nachdem wir obige Fakten als Programm in der Datei `intro.pl` gespeichert haben, können wir es am Kommando-Prompt des Interpreters einlesen:

```
| ?- [intro].
```

Falls das Programm syntaktisch korrekt ist, sollte so etwas ausgegeben werden wie:

```
compiling intro.pl for byte code...  
intro.pl compiled, 12 lines read
```

```
yes  
| ?-
```

Wenn Sie stattdessen das Programm mittels der Tastatur direkt eingeben wollen:

- Nach der Eingabe von `[user]` . kann der Benutzer Regeln und Fakten definieren.
- Mittels `Ctrl-D` wird die Eingabe beendet und automatisch der Übersetzungsvorgang gestartet.

Um jederzeit feststellen zu können, welches Programm gerade aktiv ist, geben wir folgendes ein:

```
| ?- listing.
```

Was kann man nun mit diesem Programm machen?

- Wir können fragen, ob Ruth ein Kind von Walter ist:
| `?- parent(walter, ruth).`
no
- Da Prolog annimmt, seine Fakten würden die ganze Welt komplett beschreiben, ist wahr, was gefunden wird und falsch, was nicht vorhanden ist.

Wie funktioniert das?

- Prolog vergleicht die Anfrage mit den Fakten und Regeln der Wissensbasis und prüft, ob der Funktor mit gleicher Stelligkeit vorhanden ist.
 - Der Term `walter` wird mit denen der Fakten verglichen, bei Übereinstimmung wird dasselbe mit `ruth` durchgeführt.
- Diesen Vorgang nennt man *matchen*.

Was kann man sonst noch mit diesem Programm machen?

- Wir können die Kinder von Walter auflisten:

```
| ?- parent(walter, X).
```

- Antwort:

```
X = klaus ? ; % mit dem Semikolon wird
```

```
X = eva      % die nächste Lösung angefordert
```

- Die gleiche Anfrage mit SQL auf die Tabelle Eltern:

Parent	Child
--------	-------

ruth	klaus
------	-------

walter	klaus
--------	-------

ruth	eva
------	-----

walter	eva
--------	-----

...	...
-----	-----

```
Select Child from Eltern  
where Parent = 'walter';
```

Wie funktioniert das?

- Prolog durchsucht die Fakten (von oben nach unten) und testet, ob `walter` als erster Term vorkommt (matchen).
- Anschließend wird `X` mit `klaus` *instanziiert*. (besser: belegt)
- Diesen Vorgang nennt man *Unifikation* oder *Instanziierung*. Ersetzen der Variablen durch Terme, so dass die Terme als Zeichenfolge gleich sind.
- Die gefundene Belegung wird ausgegeben.
- Die Lösungssuche endet noch nicht, da noch ungeprüfte Klauseln existieren.
- Die Lösungssuche wird bei der letzten Alternative fortgesetzt; alle Variablen, die ab diesem Punkt instanziiert wurden, werden wieder frei.
- Diesen Vorgang nennt man *Backtracking*.

Was kann man noch mit diesem Programm machen?

- Wir können die Eltern von Walter auflisten:

```
| ?- parent(X, walter).
```

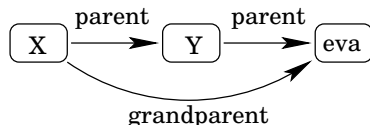
- Bedeutung der Anfrage: Für welche X lässt sich das Ziel `parent(X, walter)` aus den Fakten folgern?

- Antwort:

```
X = peter ? a    % mit dem 'a' werden alle  
X = iris         % restlichen Lösung angefordert  
no
```

Man kann „neues Wissen“ mit dem Programm generieren:

- Wir können die Großeltern von Eva ausgeben:
| `?- parent(Y, eva), parent(X, Y).`
- Zur Veranschaulichung:



- Das Komma ist als logisches UND zu lesen. Ausgabe:

```
X = hans      Y = ruth ? ;  
X = claudia   Y = ruth ? ;  
X = peter    Y = walter ? ;  
X = iris      Y = walter ? ;  
no
```

Wie funktioniert das?

- Zielliste: `parent(Y, eva), parent(X, Y)`.
 - Das Programm wird von oben nach unten durchsucht, und der passende Fakt 7 gefunden: `parent(ruth, eva)`
 - Also wird `Y=ruth` instanziiert, das erste Literal ist `true` und wird aus der Zielliste entfernt.
- Zielliste: `parent(X, ruth)`.
 - Wieder wird das Programm von oben nach unten durchsucht und Fakt 1 gefunden: `parent(hans, ruth)`
 - Also wird `X=hans` instanziiert, das erste Literal der aktuellen Zielliste ist `true` und wird aus der Zielliste entfernt.
- Die Zielliste ist leer, damit ist die Anfrage erfüllt und die Instanziierung `X=hans, Y=ruth` wird ausgegeben.

Wird eine weitere Lösung angefordert, wird die Instanziierung von X wieder rückgängig gemacht und die Suche wird vor der letzten Instanziierung fortgesetzt. → Backtracking

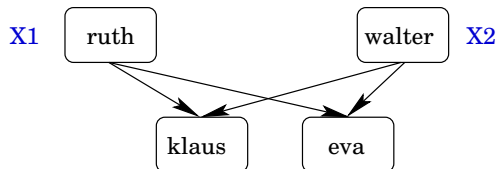
- Zielliste: `parent(X,ruth)`.
 - Im Programm wird Fakt 2 gefunden: `parent(claudia,ruth)`
 - $X=claudia$ wird instanziiert, das erste Literal der aktuellen Zielliste ist `true` und wird aus der Zielliste entfernt.
- Die Zielliste ist wieder leer, die nächste Lösung $X=claudia$, $Y=ruth$ wird ausgegeben usw.

Wie können wir testen, ob zwei Personen Geschwister sind?

- Wenn Eva und Klaus Geschwister sind, haben sie einen gemeinsamen Elternteil X:

```
| ?- parent(X, eva), parent(X, klaus).
```

- Zur Veranschaulichung:



- Ausgabe:

```
X = ruth ? ;    % mit dem Semikolon wird  
X = walter ? ; % die nächste Lösung angefordert  
no
```

Wie können wir zwischen Mutter und Vater unterscheiden?

- Zur Unterscheidung männlich/weiblich fügen wir dem Programm zwei Relationen hinzu:

<code>male(hans).</code>	<code>female(claudia).</code>
<code>male(peter).</code>	<code>female(iris).</code>
<code>male(walter).</code>	<code>female(ruth).</code>
<code>male(klaus).</code>	<code>female(celeste).</code>
<code>male(konrad).</code>	<code>female(eva).</code>
<code>male(josef).</code>	<code>female(doris).</code>

- Damit können wir z.B. Walters Vater und Mutter unterscheiden:

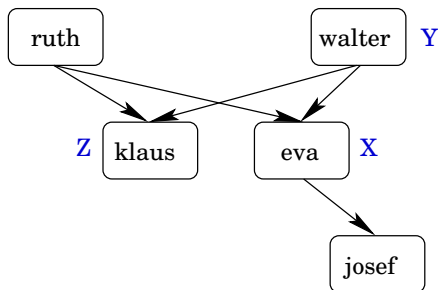
```
Vater: | ?- parent(X, walter), male(X).  
Mutter: | ?- parent(X, walter), female(X).
```

Wie können wir Onkel und Tanten von einer Person bestimmen?

- Josefs Onkel und Tante sind Geschwister seiner Eltern:

```
| ?- parent(X,josef), parent(Y,X), parent(Y,Z),  
      X \== Z.
```

- Zur Veranschaulichung:



- Hier muss zwingend X ungleich Z gelten! In Prolog: `X \== Z`

Wie können wir einem Programm *Regeln* hinzufügen?

- *Regeln:*

`mother(X,Y) :- parent(X,Y), female(X).`

`father(X,Y) :- parent(X,Y), male(X).`

- *Semantik:*

$\forall X \forall Y$ gilt: `parent(X,Y)  $\wedge$  female(X)  $\Rightarrow$  mother(X,Y)`

- *Sprechweise:*

X is the mother of Y, **if** X is a parent of Y **and** X is female.

Im Folgenden definieren wir Regeln für

- Großeltern, -vater und -mutter.
- Geschwister, Schwester, Bruder.
- Onkel und Tante.
- eine allgemeine Vorfahr-Beziehung.

Großeltern, Großvater und Großmutter:

- *Regeln:*

```
grandparent(X,Y) :- parent(X,Z), parent(Z,Y).  
grandfather(X,Y) :- grandparent(X,Y), male(X).  
grandmother(X,Y) :- grandparent(X,Y), female(X).
```

- *Semantik:* $\forall X \forall Y \forall Z$ gilt:

$\text{parent}(X,Z) \wedge \text{parent}(Z,Y) \Rightarrow \text{grandparent}(X,Y)$

- *Sprechweise:*

X is grandparent of Y, if X is parent of Z and Z is parent of Y.

- Beispiel:

```
| ?- grandparent(X,josef).  
X = ruth ? ;  
X = walter ? ;  
no
```

Wie funktioniert das?

- Zielliste: `grandparent(X, josef)`.
 - Das Programm wird von oben nach unten durchsucht und die erste passende Regel gefunden:
`grandparent(X,Y) :- parent(X,Z), parent(Z,Y)`.
 - Also wird `Y=josef` instanziiert und der Eintrag der Zielliste wird ersetzt durch die rechte Seite der Regel.
- Zielliste: `parent(X,Z), parent(Z,josef)`.
 - Es wird eine passende Regel für den ersten Eintrag der Zielliste gesucht und die Regel `parent(hans,ruth)` gefunden.
 - Also wird `X=hans` und `Z=ruth` instanziiert und der erste Eintrag der Zielliste durch `true` ersetzt, also entfernt.
- Zielliste: `parent(ruth,josef)`.
 - Dieser Fakt wird nicht im Programm gefunden, das Ziel ist mit dieser Instanziierung nicht erfüllt.
 - Die letzte Belegung wird zurückgenommen und als Zielliste ergibt sich wieder: `parent(X,Z), parent(Z,josef)`.

Wie funktioniert das? (Fortsetzung)

- Zielliste: `parent(X,Z)` , `parent(Z,josef)` .
 - Es wird der nächste Fakt `parent(claudia, ruth)` untersucht und `X=claudia` und `Z=ruth` instanziiert.
 - Der erste Eintrag der Zielliste wird durch `true` ersetzt, also entfernt.
- Zielliste: `parent(ruth,josef)` .
 - Wieder wird kein übereinstimmender Fakt gefunden und ein Backtracking durchgeführt.
- Zielliste: `parent(X,Z)` , `parent(Z,josef)` .
 - Dieser Vorgang wiederholt sich so lange, bis der Fakt `parent(ruth,eva)` gefunden wird, also `X=ruth` und `Z=eva`.
- Zielliste: `parent(eva,josef)`
 - Dieses Ziel wird als Fakt gefunden und die Instanziierung `X=ruth` ausgegeben. Der Wert von `Z` wird nicht ausgegeben, da er nicht in der ursprünglichen Zielliste auftaucht.

Geschwister haben einen gemeinsamen Elternteil:

- *Regeln:*

```
sibling(X,Y) :- parent(Z,X), parent(Z,Y),  
                X \== Y.
```

```
sister(X,Y) :- sibling(X,Y), female(X).
```

```
brother(X,Y) :- sibling(X,Y), male(X).
```

- *Semantik:* $\forall X \forall Y \forall Z$ gilt:

```
parent(Z,X)  $\wedge$  parent(Z,Y)  $\wedge$   $X \neq Y \Rightarrow$   
sibling(X,Y)
```

- *Sprechweise:*

X is sibling of Y, if Z is parent of X and Z is parent of Y and X is not equal to Y.

- Onkel und Tante

```
aunt(X,Y) :- sibling(X,Z), parent(Z,Y),  
            female(X).
```

```
uncle(X,Y) :- sibling(X,Z), parent(Z,Y),  
             male(X).
```

- eine allgemeine Vorfahr-Beziehung *mittels Rekursion*

```
ancestor(X,Y) :- parent(X,Y).
```

```
ancestor(X,Y) :- parent(X,Z), ancestor(Z,Y).
```

- Allgemeines
- Erste Schritte
- *Arithmetik*
- Listen
- Backtracking
- Sortieren
- Input/Output
- Zahlenrätsel
- Kannibalen und Missionare
- 8-Damen-Problem
- Schiebepuzzle
- A^* -Algorithmus

Die bekannten arithmetischen Operationen sind definiert:

- + Addition
- Subtraktion
- * Multiplikation
- / Division (floating point)
- // Division (integer)
- mod Rest bei ganzzahliger Division

Allerdings wird die Anfrage

```
| ?- X = 1 + 2.
```

nicht mit $X = 3$, sondern mit $X = 1 + 2$ beantwortet!

Der Ausdruck $1 + 2$ ist in Prolog ein Term, wobei $+$ die Funktion und 1 und 2 die Operanden sind. Die Addition wird nur dann durchgeführt, wenn der Operator `is` verwendet wird:

```
| ?- X is 1 + 2.
```

Werden numerische Werte verglichen, erfolgt auch automatisch eine Auswertung der arithmetischen Ausdrücke:

| ?- 7 * 3 > 20.
yes

| ?- 7 * 3 < 20.
no

Dies gilt für die folgenden Operatoren:

$X > Y$	X is greater than Y
$X < Y$	X is less than Y
$X \geq Y$	X is greater than or equal to Y
$X \leq Y$	X is less than or equal to Y
$X == Y$	the values of X and Y are equal
$X \neq Y$	the values of X and Y are not equal

Beachten Sie den Unterschied zwischen $=$ und $==$ sowie den Unterschied zwischen $\neq$ und $\neq$.

Wichtig: Machen Sie sich den Unterschied zwischen den beiden Operatoren `=` und `==` klar!

- Bei `X = Y` versucht Prolog, die Objekte `X` und `Y` zu unifizieren (gleich zu machen) und ggf. mit Werten zu belegen.

?- 1 + 2 = 2 + 1.	?- 1 + X = Y + 3.
no	X = 3
	Y = 1
	yes

- Der Operator `==` wertet den arithmetischen Ausdruck aus, belegt die Variablen aber nicht mit Werten.

?- 1 + 2 == 2 + 1.	?- 1 + X == Y + 3.
yes	uncaught exception:

Im zweiten Beispiel hätten `X` und `Y` vorher mit Werten belegt werden müssen, z.B. `X=4`, `Y=2`, `1+X == Y+3`.

Test auf Gleichheit:

= unifizierbar
== identisch mit
:= numerisch gleich

Test auf Ungleichheit:

\= nicht unifizierbar
\== nicht identisch mit
=\= numerisch ungleich

Beispiel:

```
| ?- X=1+3, Y=2+2, X:=Y.  
yes
```

Beim Vergleich $X:=Y$ werden die Terme $1+3$ und $2+2$ zunächst ausgewertet und dann verglichen.

```
| ?- X=1+3, Y=2+2, X==Y.  
no
```

Bei dem Vergleich $X==Y$ wird geprüft, ob die Inhalte von X und Y identisch sind (ohne eine arithmetische Auswertung).

Beispiel:

```
| ?- 1+X = 1+Y.
```

```
Y = X
```

```
yes
```

Die Variablen X und Y können unifiziert werden. Hier sieht man, das Prolog einen allgemeinsten Unifikator sucht. Es werden nicht alle Möglichkeiten aufgezählt wie $X=1, Y=1$ und $X=2, Y=2$ usw., stattdessen wird nur die „minimale“ Unifikation erzeugt.

```
| ?- 1+X == 1+Y.
```

```
no
```

Prolog versucht nicht, die Terme gleich zu machen, es findet hier keine Unifikation statt, es werden lediglich zwei Zeichenketten miteinander verglichen.

Übung: Schreiben Sie ein Prädikat `maximum(X,Y,Z,Erg)`, das den maximalen Wert der drei Variablen X, Y und Z bestimmt und in Erg bereit stellt.

Übung: Schreiben Sie ein Prädikat `fib(N,Erg)`, das die N-te Fibonacci-Zahl F_N in Erg berechnet.

Übung: Schreiben Sie ein Prädikat `ggT(X,Y,Erg)` zur Berechnung des größten gemeinsamen Teilers zweier ganzer Zahlen X und Y.

Zur Erinnerung:

$$\text{ggT}(x, y) = \begin{cases} x, & \text{falls } y = 0 \\ \text{ggT}(y, x \bmod y), & \text{sonst} \end{cases}$$

In Prolog können wir zur Laufzeit Fakten und Regeln der Datenbasis hinzufügen.

- `asserta(term)` fügt `term` am Anfang der Datenbasis ein.
- `assertz(term)` fügt `term` am Ende der Datenbasis an.
- `retract(term)` entfernt `term` wieder aus der Datenbasis.

Damit das funktioniert, muss der Term entweder unbekannt sein, oder das Prädikat muss als `dynamic` vereinbart sein.

Um die Laufzeit des `fib`-Prädikats zur Berechnung der Fibonacci-Zahlen zu verbessern, vermeiden wir rekursive Aufrufe durch memorieren bereits berechneter Werte.

```
:- dynamic(fib/2).
```

```
fib(0,0).
```

```
fib(1,1).
```

```
fib(N, Erg) :-
```

```
    N > 0,
```

```
    N1 is N-1,
```

```
    N2 is N-2,
```

```
    fib(N1,E1),
```

```
    fib(N2,E2),
```

```
    Erg is E1+E2,
```

```
    asserta(fib(N, Erg)).
```

Vergleichen Sie die Laufzeiten dieser und der zuvor erstellten Lösung. Wie groß darf die Eingabe N werden?

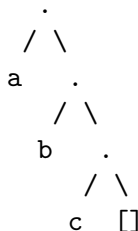
- Allgemeines
- Erste Schritte
- Arithmetik
- *Listen*
- Backtracking
- Sortieren
- Input/Output
- Zahlenrätsel
- Kannibalen und Missionare
- 8-Damen-Problem
- Schiebepuzzle
- A^* -Algorithmus

Eine Liste ist eine Aufzählung von Elementen und wird in Prolog in eckigen Klammern [...] geschrieben:

```
[adam, eva, heidi, peter, bonnie, clyde]  
[0, 1, 2, 3, 4, 5, 6, 7, 8, 9]
```

Wir können eine Liste [a,b,c] auch in Paar-Schreibweise angeben: `.(a,.(b,.(c,[])))`

Listen werden in Prolog als entartete Binärbäume gespeichert:



Daraus ergibt sich die dritte Darstellung: Aufzählung mit Restliste

- Eine leere Liste wird geschrieben als `[]`.
- Eine nicht-leere Liste besteht immer
 - aus dem ersten Element, Kopf (engl.: head) der Liste genannt,
 - und dem Rest, Schwanz (engl.: tail) der Liste genannt.

Der senkrechte Strich teilt den Kopf vom Schwanz der Liste:

$$L = [a, b, c, d]$$
$$L = [a \mid \text{Tail}]$$

Unifizieren wir `[X|Rest]` mit `[1,2,3,4]`, dann erhalten wir:

$$X = 1$$
$$\text{Rest} = [2,3,4]$$

Listen können auch ineinander verschachtelt werden:

```
| ?- Hobbies1 = [rennen, tanzen],  
    Hobbies2 = [schlafen, essen],  
    L = [heidi, Hobbies1, peter, Hobbies2].
```

```
Hobbies1 = [rennen,tanzen]  
Hobbies2 = [schlafen,essen]  
L = [heidi,[rennen,tanzen],peter,[schlafen,essen]]
```

Wir betrachten im Weiteren folgende Operationen:

- `member(X, L)` testet, ob das Objekt `X` in der Liste `L` enthalten ist.
- `sublist(L1, L)` testet, ob `L1` eine Teilliste von `L` ist.
- `permutation(L, P)` erzeugt in `P` alle Permutationen der in `L` enthaltenen Objekte.

`member(X, L)` testet, ob das Objekt `X` in der Liste `L` enthalten ist:

```
| ?- member(b, [a,b,c]).
```

```
true
```

```
| ?- member(b, [a,[b,c],d]).
```

```
no
```

```
| ?- member([b,c], [a,[b,c],d]).
```

```
true
```

`X` ist Element der Liste `L`, wenn

- `X` der Kopf der Liste ist, also

```
member(X, [X|Tail]).
```

- oder `X` im Rest der Liste enthalten ist, also

```
member(X, [Head|Tail]) :- member(X, Tail).
```

Übung: Schreiben Sie eine Funktion `conc(L1, L2, L)`, die die Listen L1 und L2 aneinander hängt. In der Ergebnisliste L stehen also zunächst alle Elemente von L1, dann alle Elemente von L2.

Übung: Schreiben Sie eine Funktion `add(X,L1,L)`, die ein Element X am Anfang der Liste L1 einfügt, Liste L ist die Ergebnisliste.

Übung: Schreiben Sie eine Funktion `del(X,L1,L)`, die ein Element X aus der Liste L1 entfernt und Liste L liefert.

Was liefert der Aufruf `del(1, L, [2,3,4])`?

Übung: Schreiben Sie eine Funktion `insert(X,L1,L)`, die ein Element X an einer beliebigen Stelle der Liste L1 einfügt und Liste L liefert.

Debugging: Mittels trace aktivieren Sie den Debugger, mittels notrace wird der Debugger ausgeschaltet.

```
| ?- trace.  
| ?- insert(1, [4,3,2], L).  
      1      1  Call: insert(1,[4,3,2],_22) ?  
      1      1  Exit: insert(1,[4,3,2],[1,4,3,2]) ?  
  
L = [1,4,3,2] ? ;  
      1      1  Redo: insert(1,[4,3,2],[1,4,3,2]) ?  
      2      2  Call: insert(1,[3,2],_55) ?  
      2      2  Exit: insert(1,[3,2],[1,3,2]) ?  
      1      1  Exit: insert(1,[4,3,2],[4,1,3,2]) ?  
| ?- notrace.
```

`sublist(L1, L)` testet, ob L1 eine Teilliste von L ist:

```
| ?- sublist([c,d,e], [a,b,c,d,e,f]).  
true
```

```
| ?- sublist([c,e], [a,b,c,d,e,f]).  
no
```

S ist eine Teilliste von L, wenn

- L in zwei Listen L1 und L2 aufgeteilt werden kann und
- die Liste L2 wiederum in die Liste S und eine Liste L3 aufgeteilt werden kann.

```
sublist(S,L) :- conc(L1,L2,L), conc(S,L3,L2).
```

- Stellen Sie diesen Sachverhalt graphisch dar.
- Was liefert der Aufruf `sublist(S, [1,2,3]).?`

Übung: Die in gprolog eingebaute Funktion `sublist` verhält sich etwas anders:

```
| ?- sublist([c,e], [a,b,c,d,e,f]).  
true
```

```
| ?- sublist([e,c], [a,b,c,d,e,f]).  
no
```

Wie könnte diese Funktion implementiert sein?

`permutation(L, P)` erzeugt in `P` alle Permutationen der in `L` enthaltenen Objekte

```
| ?- permutation([a,b,c], P).  
P = [a,b,c];  
P = [a,c,b];  
P = [b,a,c];  
....
```

Mögliche Implementierung:

- Abbruchbedingung: `permutation([], []).`
- Falls die erste Liste die Form `[X|L]` hat: Erstelle zunächst Liste `L1` als Permutation von `L`, füge dann `X` an jeder Stelle in `L1` ein.

```
permutation([X|L], P) :- permutation(L, L1),  
                           insert(X,L1,P).
```

Übung: Geben Sie eine mögliche Implementierung der in gprolog eingebauten Relation `reverse` an.

Beispiel: `reverse([a,b,c,d], P)` liefert `P = [d,c,b,a]`

Übung: Definieren Sie eine Relation `shift(L1, L2)`, so dass L2 aus L1 hervorgeht, indem die Elemente zyklisch nach links verschoben werden.

Beispiel:

```
| ?- shift([a,b,c,d,e], L1),  
      shift(L1, L2).
```

liefert

`L1 = [b,c,d,e,a]`

`L2 = [c,d,e,a,b]`

Übung: Schreiben Sie ein Prädikat `lastElem(L,X)`, das das letzte Listenelement der Liste L in der Variablen X liefert.

Übung: Schreiben Sie Prädikate, mit denen man feststellen kann, ob eine Liste Präfix (bzw. Suffix) einer anderen Liste ist.

Anonyme Variablen: Häufig benötigt man einen Wert nicht weiter. In einem solchen Fall kann man eine anonyme Variable verwenden, die mit einem Unterstrich beginnt. Ihr Wert wird nicht gespeichert:

`erstesElement(X, [X|_Rest]).`

- Allgemeines
- Erste Schritte
- Arithmetik
- Listen
- *Backtracking*
- Sortieren
- Input/Output
- Zahlenrätsel
- Kannibalen und Missionare
- 8-Damen-Problem
- Schiebepuzzle
- A^* -Algorithmus

Damit Prolog für praktische Probleme zugänglicher wird, gibt es den Cut-Operator '!', mit dem das Backtracking durch Teilbäume des Regelwerks unterbunden werden kann.

Die Lösungsmenge ist dann zwar nicht mehr in jedem Fall identisch, aber Programme können unter Performance-Aspekten optimiert werden.

Steht in einer Klausel ein Cut und wird diese Stelle bei der Auswertung erreicht, so werden alle Belegungen, die beim Matching der vorangehenden Prädikate getroffen wurden eingefroren. Für sie findet kein Backtracking mehr statt, es werden auch keine Klauseln der gleichen Prozedur mehr ausprobiert.

Scheitert der Regelteil hinter dem Cut, so scheitert die gesamte Prozedur.

Der Cut: grüne und rote Cuts

Ein Cut kann eingesetzt werden, um die Berechnung zu beschleunigen, ohne dass Lösungen verlorengehen oder hinzukommen. Solche Cuts heißen *grüne Cuts*.

```
minimum(X,Y,X) :- X < Y, !.  
minimum(X,Y,Y) :- X >= Y, !.
```

Der zweite Cut kann weggelassen werden, da die Prozedur `minimum` dann abgeschlossen ist. (Eine Zusammenfassung von „Regeln gleichen Namens“ nennt man auch Prozedur.)

Ohne Cuts hat das Programm die gleiche Lösungsmenge.

Der Cut: grüne und rote Cuts

Alternativ: `minimum(X,Y,Z) :- X < Y, !, Z=X.`
 `minimum(X,Y,Y).`

Hier berücksichtigt man die SLD-Strategie: Auswerten der Klauseln von oben nach unten und innerhalb einer Klausel von links nach rechts.

- Aufruf `minimum(2,3,3)` liefert bei Verwendung des Cuts `no`.
 - Zunächst wird `X < Y` als `true` bewertet, der Cut wird als `true` interpretiert, die Unifikation `Z=X` schlägt fehl.
 - Deshalb geht der Prolog-Interpreter die Literale der Klausel zurück, um ein Backtracking zu ermöglichen.
 - Dabei trifft er auf den Cut, das Backtracking wird - auf dieser Ebene - abgebrochen, damit ist insgesamt die Auswertung fehlgeschlagen.
- Der Aufruf `minimum(2,3,3)` liefert ohne Verwendung des Cuts `yes`.
- Dies ist ein *roter Cut*.

- Allgemeines
- Erste Schritte
- Arithmetik
- Listen
- Backtracking
- *Sortieren*
- Input/Output
- Zahlenrätsel
- Kannibalen und Missionare
- 8-Damen-Problem
- Schiebepuzzle
- A*-Algorithmus

Idee:

- Tausche zwei benachbarte Zahlen X und Y , falls $X > Y$.
- Sortiere den Rest.

```
% Wenn ein Wertepaar getauscht wurde, also swap  
% erfolgreich war, darf die nächste Regel NICHT  
% ausgeführt werden! Roter Cut!
```

```
bubblesort(List, SortedList) :-  
    swap(List, List2), !,  
    bubblesort(List2, SortedList).
```

```
bubblesort(Sorted, Sorted).
```

```
swap([X,Y|Rest], [Y,X|Rest]) :- X > Y.  
swap([Z|Rest], [Z|Rest2]) :- swap(Rest, Rest2).
```

Idee: Sortiere nicht-leere Liste $L=[X|Tail]$ wie folgt:

- Sortiere den Rest *Tail* der Liste *L* zu *SortedTail*.
- Füge *X* an der richtigen Stelle in *SortedTail* ein.

```
insertionsort([], []).
```

```
insertionsort([X|Tail], SortedList) :-  
    insertionsort(Tail, SortedTail),  
    insert(X, SortedTail, SortedList).
```

% Wenn diese Regel angewendet wird, darf die
% nächste Regel NICHT angewendet werden. Roter Cut!

```
insert(X, [Y|Sorted], [Y|Sorted2]) :-  
    X > Y, !,  
    insert(X, Sorted, Sorted2).  
insert(X, Sorted, [X|Sorted]).
```

Idee: Sortiere eine Liste L wie folgt:

- Wähle das erste Element X der Liste L als Pivot-Element.
- Teile mittels `split` die Restliste auf in eine Liste mit kleineren Werten `Small` und eine Liste mit größeren Werten `Big` als X.
- Sortiere `Small` rekursiv zu `SortedSmall`.
- Sortiere `Big` rekursiv zu `SortedBig`.
- Füge die Listen mittels `conc` zusammen zur sortierten Liste `SortedList`.

```
quicksort([], []).
```

```
quicksort([X|Tail], SortedList) :-  
    split(X, Tail, Small, Big),  
    quicksort(Small, SortedSmall),  
    quicksort(Big, SortedBig),  
    conc(SortedSmall, [X|SortedBig], SortedList).
```

Quicksort

```
%%% helper functions for quicksort  
split(X, [], [], []).
```

```
split(X, [Y|Tail], [Y|SmallItems], BigItems) :-  
    X > Y, !,      % Roter Cut!  
    split(X, Tail, SmallItems, BigItems).
```

```
split(X, [Y|Tail], SmallItems, [Y|BigItems]) :-  
    split(X, Tail, SmallItems, BigItems).
```

```
conc([], L, L).  
conc([X|L1], L2, [X|L3]) :- conc(L1, L2, L3).
```

- Allgemeines
- Erste Schritte
- Arithmetik
- Listen
- Backtracking
- Sortieren
- *Input/Output*
- Zahlenrätsel
- Kannibalen und Missionare
- 8-Damen-Problem
- Schiebepuzzle
- A^* -Algorithmus

In Prolog gibt es build-in predicates zur Interaktion mit dem Anwender:

- `write(X)` gibt den Inhalt des Terms `X` auf dem Bildschirm aus.
 - `nl` steht für new line.
 - `tab(N)` fügt `N` Leerzeichen in die Ausgabe ein.
- `read(X)` liest den nächsten Term von der Tastatur und belegt `X` mit dem gelesenen Wert.

Die Eingabe muss mit einem Punkt abgeschlossen werden.

Beispiel:

```
cube :-  
    write('next value, please: '),  
    read(N),  
    C is N*N*N,  
    write(C).
```

Übung: Schreiben Sie ein Prädikat `writelist(L)`, das die einzelnen Listen einer Liste von Listen `L` jeweils in einer Zeile ausgibt.

Ausgabe von `writelist([[1,2,3], [3,4,5,6]])`:

1,2,3

3,4,5,6

Übung: Schreiben Sie ein Prädikat `bars(L)`, das zu jeder Zahl der Liste jeweils eine Zeile mit entsprechend vielen Sternchen ausgibt.

Ausgabe von `bars([3,6,5])`:

Dateioperationen:

- `tell(File)` leitet die Ausgabe der Prädikate `write`, `put`, `nl` usw. in die Datei `File` um.
- `told` beendet die Umleitung und bewirkt, dass spätere Ausgaben wieder auf dem Bildschirm erscheinen.
- `telling(File)` gibt an, in welche Datei die Ausgabe zur Zeit erfolgt, dabei steht `user` für die Tastatur.

```
| ?- telling(Anfang),  
    tell('datei.txt'),      % öffnen der Datei  
    telling(Mitte),  
    write('Hallo, '), nl,   % erste Zeile schreiben  
    write('Welt!'), nl,    % zweite Zeile schreiben  
    told,                  % schließen der Datei  
    telling(Ende).
```

Dateioperationen:

- `see(File)` nimmt die Eingabe für die Prädikate `read`, `get` usw. aus der Datei `File`.
- `seen` beendet die Umleitung und bewirkt, dass spätere Eingaben wieder vom Benutzer abgefragt werden.
- `seeing(File)` gibt an, aus welcher Datei die Eingabe zur Zeit genommen wird, dabei steht `user` für die Tastatur.

```
| ?- seeing(Anfang),  
    see('datei.txt'),      % öffnen der Datei  
    seeing(Mitte),  
    get0(X), put(X),       % erstes Zeichen lesen  
    get0(Y), put(Y),      % zweites Zeichen lesen  
    seen,                 % schließen der Datei  
    seeing(Ende).
```

Das Prädikat `read` kann nur dann zum Lesen verwendet werden, wenn Terme gelesen werden, also bspw. Zahlen, Fakten oder Regeln, die allesamt mit einem Punkt abgeschlossen sind.

```
test :-  
    tell('datei.txt'),      % öffnen der Datei  
    write('sunshine.'), nl,  
    write('raining.'), nl,  
    write('funny :- sunshine, raining.'), nl,  
    write('nice :- sunshine, \\+raining.'), nl,  
    write('.'),  
    told.                  % schließen der Datei
```

In GNU-Prolog wird `not` als `\\+` dargestellt und muss innerhalb einer Zeichenkette natürlich maskiert werden.

Die in einer Datei abgelegten Fakten und Regeln können wie folgt eingelesen und der Datenbasis hinzugefügt werden:

```
add_to_database(File) :-  
    see(File),          /* open file */  
    repeat,  
    read(Data),         /* read from File */  
    process(Data),  
    seen,               /* close File */  
    !.                 /* stop reading */  
  
process(.) :- !.  
process(Data) :- asserta(Data), nl, fail.
```

Das Prädikat **repeat** ist bei jeder Auswertung wahr und kann verwendet werden, um Schleifen zu programmieren.

- Allgemeines
- Erste Schritte
- Arithmetik
- Listen
- Backtracking
- Sortieren
- Input/Output
- *Zahlenrätsel*
- Kannibalen und Missionare
- 8-Damen-Problem
- Schiebepuzzle
- A^* -Algorithmus

Wir wollen zunächst ein Rätsel der folgenden Form lösen:

			MEHR	
SEND	AAL	BILL	+ MATHE	GERALD
+ MORE	+ AAL	+ IRMA	+ MACHT	+ DONALD
-----	-----	-----	-----	-----
MONEY	FANG	LIEBE	FREUDE	ROBERT

Eine naive Implementierung ohne Listen könnte zunächst die einzelnen Ziffern definieren:

digit(1).	digit(6).
digit(2).	digit(7).
digit(3).	digit(8).
digit(4).	digit(9).
digit(5).	digit(0).

Die einzelnen Ziffern müssen paarweise verschieden sein, die ersten Ziffern müssen ungleich 0 sein.

```
solve :-                                     % SEND + MORE = MONEY
    digit(S), S\=0,
    digit(E), E\=S,
    digit(N), N\=E, N\=S,
    digit(D), D\=N, D\=E, D\=S,
    digit(M), M\=D, M\=N, M\=E, M\=S, M\=0,
    digit(O), O\=M, O\=D, O\=N, O\=E, O\=S,
    digit(R), R\=O, R\=M, R\=D, R\=N, R\=E, R\=S,
    digit(Y), Y\=R, Y\=O, Y\=M, Y\=D, Y\=N, Y\=E, Y\=S,
    sum(S,E,N,D,M,O,R,Y),
    write('S = '), write(S), nl,
    ....
    write('Y = '), write(Y), nl.
```

Abschließend müssen wir noch prüfen, ob die gegebene Bedingung $\text{SEND} + \text{MORE} = \text{MONEY}$ ist:

```
sum(S,E,N,D,M,O,R,Y) :-
```

```
    N1 is      1000*S + 100*E + 10*N + D,
```

```
    N2 is      1000*M + 100*O + 10*R + E,
```

```
    N3 is 10000*M + 1000*O + 100*N + 10*E + Y,
```

```
    N3 is N1 + N2.
```

Fragen:

- Wie können wir die Laufzeit des Programms verbessern?
- Können wir Listen zur Lösung des Problems nutzen, um das Programm kürzer zu gestalten?

Naives Generate-And-Test-Verfahren:

- Erzeuge eine Permutation für die Belegung der Variablen,
`gen(S,E,N,D,M,O,R,Y) :-
 permutation([S,E,N,D,M,O,R,Y,_,_],
 [0,1,2,3,4,5,6,7,8,9]).`
- teste, ob die gegebene Gleichung für die Belegung erfüllt ist
`gl(S,E,N,D,M,O,R,Y) :- M \= 0, S \= 0,
 (S*1000 + E*100 + N*10 + D) +
 (M*1000 + O*100 + R*10 + E) =:=
 (M*10000 + O*1000 + N*100 + E*10 + Y).`
- und gib die Belegung aus.
`solve :- gen(S,E,N,D,M,O,R,Y),
 gl(S,E,N,D,M,O,R,Y),
 write('S = '), write(S), nl,
 write('E = '), write(E), nl,`

Wir können es natürlich auch einmal mit dem Einsatz von Wissen versuchen! Hier können wir Wissen über die Addition einbringen:

```
  S  E  N  D
+ M  O  R  E
+ C3 C2 C1
-----
M  O  N  E  Y
```

```
10 * C1 + Y == D + E
10 * C2 + E == N + R + C1
10 * C3 + N == E + O + C2
10 *  M + O == S + M + C3
M = 1      (Übertrag)
```

Wir formen die obigen Gleichungen um und prüfen nach jeder Belegung einer Variablen, ob die Teilbelegung noch richtig ist:

solve :-

```
M = 1,                % M muss 1 sein
L = [2,3,4,5,6,7,8,9,0], % restliche Ziffern
select(S, L, L1),      % wähle Ziffer für S
S > 0,                 % S muss ungleich 0 sein
```

$(C3 = 0; C3 = 1),$	% Übertrag ist 0 oder 1
$0 \text{ is } S + M + C3 - 10 * M,$	% 0 ist jetzt eindeutig
$\text{select}(0, L1, L2),$	% entferne 0 aus Rest
$\text{select}(E, L2, L3),$	% wähle Ziffer für E
$(C2 = 0; C2 = 1),$	% Übertrag ist 0 oder 1
$N \text{ is } E + 0 + C2 - 10 * C3,$	% N ist jetzt eindeutig
$\text{select}(N, L3, L4),$	% entferne N aus Rest
$(C1 = 0; C1 = 1),$	% Übertrag ist 0 oder 1
$R \text{ is } E + 10 * C2 - N - C1,$	% R ist jetzt eindeutig
$\text{select}(R, L4, L5),$	% entferne R aus Rest

```
select(D, L5, L6),          % wähle Ziffer für D
Y is D + E - 10*C1,         % Y ist jetzt eindeutig
select(Y, L6, _),          % entferne Y aus Rest

% und dann noch die Ausgabe
write('S = '), write(S), nl,
write('E = '), write(E), nl,
....
write('Y = '), write(Y), nl.
```

Laufzeitvergleich?

Wir wollen nun ein Programm schreiben, das obige Idee aufgreift, aber Zahlenrätsel mittels Listen löst.

Wenn wir zwei Zahlen addieren, gehen wir spaltenweise von rechts nach links vor. In einer Spalte addieren wir einen Übertrag C1 von rechts und zwei Ziffern D1 und D2. Wir erhalten als Ergebnis eine Ziffer D und einen Übertrag C nach links.

			D1			
			D2			
			C1		0	

0			C			
			D			

Das Carry-Bit in der rechten Spalte ist 0, ebenso darf links vom Ergebnis kein Übertrag-Bit auftreten. Außerdem gilt:

- $D = (D1 + D2 + C1) \bmod 10$
- $C = (D1 + D2 + C1) \div 10$

Wir definieren ein Prädikat `sum(N1, N2, N)`, das die zwei Zahlen N1 und N2, repräsentiert als Liste, addiert und das Ergebnis als Liste N speichert: `sum([0,A,A,L], [0,A,A,L], [F,A,N,G])`.

Das verwendete Prädikat `sum1(N1,N2,N,C1,C2,Digs1,Digs2)` wertet zusätzliche Informationen aus:

- `C1`: Übertrag-Bit vor der Summation
- `C2`: Übertrag-Bit nach der Summation
- `Digs1`: Menge der verfügbaren Ziffern vor der Summation
- `Digs2`: Menge der verbleibenden Ziffern nach der Summation

```
sum(N1, N2, N) :-  
    sum1(N1, N2, N, 0, 0, [0,1,2,3,4,5,6,7,8,9], _).
```

```
sum1([], [], [], 0, 0, Digits, Digits).  
sum1([D1|N1], [D2|N2], [D|N], C1, C, Digs1, Digs) :-  
    sum1(N1, N2, N, C1, C2, Digs1, Digs2),  
    digitsum(D1, D2, C2, D, C, Digs2, Digs).
```

Das Prädikat `digitsum(D1,D2,C1,D,C,Digs1,Digs)`

- wählt aus den noch verfügbaren Ziffern in `Digs1` die Ziffern `D1`, `D2` und `D` aus, so dass nach der Summation der Spalte noch die Ziffern in `Digs` zur Verfügung stehen.
- berechnet den Wert $D = (D1+D2+C1) \bmod 10$ sowie den Wert $C = (D1+D2+C1) \div 10$.

```
digitsum(D1, D2, C1, D, C, Digs1, Digs) :-  
    del(D1, Digs1, Digs2),  
    del(D2, Digs2, Digs3),  
    del(D, Digs3, Digs),  
    S is D1 + D2 + C1,  
    D is S mod 10,  
    C is S // 10.
```

Wir müssen darauf achten, dass eine bereits instanziierte Variable D1, D2 oder D nicht nochmals aus der Liste Digs1 entfernt wird.

```
% do nothing if A is already instantiated
del(A, L, L) :- nonvar(A), !.
del(A, [A|L], L).
del(A, [B|L], [B|L1]) :- del(A, L, L1).

puzzle1([O,S,E,N,D], [O,M,O,R,E], [M,O,N,E,Y]).
puzzle2([O,B,I,L,L], [O,I,R,M,A], [L,I,E,B,E]).
puzzle3([D,O,N,A,L,D], [G,E,R,A,L,D], [R,O,B,E,R,T]).
puzzle4([O,A,A,L], [O,A,A,L], [F,A,N,G]).
```

Aufrufbeispiel: `puzzle1(N1,N2,N3), sum(N1,N2,N3).`

Es gibt natürlich auch kompliziertere Zahlenrätsel:

$$ABB - CD = EED$$

$$\quad - \quad - \quad *$$

$$FD + EF = CE$$

$$= \quad = \quad =$$

$$EGD * FH = \text{????}$$

$$ACE + DAC = JFD$$

$$\quad - \quad + \quad -$$

$$AAA - HFC = GI$$

$$= \quad = \quad =$$

$$AH + III = JBJ$$

Ein naives Generate-And-Test-Verfahren testet einfach wieder alle möglichen Belegungen:

```
gen(A,B,C,D,E,F,G,H) :-  
    permutation([A,B,C,D,E,F,G,H,_,_],  
                [0,1,2,3,4,5,6,7,8,9]).
```

Anschließend testen wir, ob alle Gleichungen erfüllt sind:

```
gl1(A,B,C,D,E) :- (E*100 + E*10 + D) ==  
    ((A*100 + B*10 + B) - (C*10 + D)).
```

```
gl2(C,D,E,F) :- (C*10 + E) ==  
    ((F*10 + D) + (E*10 + F)).
```

```
gl3(A,B,D,E,F,G) :- (E*100 + G*10 + D) ==  
    ((A*100 + B*10 + B) - (F*10 + D)).
```

```
gl4(C,D,E,F,H) :- (F*10 + H) ==  
    ((C*10 + D) - (E*10 + F)).
```

```
gl5(C,D,E,F,G,H,X) :-  
    ((E*100 + E*10 + D) * (C*10 + E)) ==  
    ((E*100 + G*10 + D) * (F*10 + H)),  
    X is ((E*100 + G*10 + D) * (F*10 + H)).
```

Und nun alles zusammen packen:

```
solve :- gen(A,B,C,D,E,F,G,H),  
         gl1(A,B,C,D,E),  
         gl2(C,D,E,F),  
         gl3(A,B,D,E,F,G),  
         gl4(C,D,E,F,H),  
         gl5(C,D,E,F,G,H,X),  
         write('A = '), write(A), nl,  
         write('B = '), write(B), nl,  
         ....  
         write('H = '), write(H), nl,  
         write('X = '), write(X).
```

- Allgemeines
- Erste Schritte
- Arithmetik
- Listen
- Backtracking
- Sortieren
- Input/Output
- Zahlenrätsel
- *Kannibalen und Missionare*
- 8-Damen-Problem
- Schiebepuzzle
- A*-Algorithmus

Erinnern wir uns:

- 3 Missionare und 3 Kannibalen wollen einen Fluss überqueren.
- Es steht nur ein Ruderboot für die Überfahrt zur Verfügung.
- Das Boot kann maximal 2 Personen aufnehmen.
- Es dürfen nie mehr Kannibalen als Missionare an einem Ort sein, sonst gibt es ein großes Fressen.

Zunächst definieren wir einen gültigen Zustand:

```
% keine Missionare --> beliebig viele Kannibalen
zust( 0, K1, _, _, Kr) :- K1 >= 0, Kr >= 0.
zust( _, K1, _, 0, Kr) :- K1 >= 0, Kr >= 0.
% sonst: Anzahl Missionare >= Anzahl Kannibalen
zust(M1, K1, _, Mr, Kr) :- M1 >= K1, Mr >= Kr,
                           K1 >= 0, Kr >= 0.
```

Dann definieren wir die Überfahrten vom linken zum rechten Ufer:

```
% 1 Kannibale von Links nach Rechts
zug(n(Ml,Kl,'L',Mr,Kr), n(Ml,Kl2,'R',Mr,Kr2)) :-
    Kl2 is Kl - 1,
    Kr2 is Kr + 1,
    zust(Ml,Kl2,'R',Mr,Kr2).
```

```
% 2 Kannibalen von Links nach Rechts
zug(n(Ml,Kl,'L',Mr,Kr), n(Ml,Kl2,'R',Mr,Kr2)) :-
    Kl2 is Kl - 2,
    Kr2 is Kr + 2,
    zust(Ml,Kl2,'R',Mr,Kr2).
```

```
% 1 Kannibale, 1 Missionar von Links nach Rechts  
zug(n(M1,K1,'L',Mr,Kr), n(M12,K12,'R',Mr2,Kr2)) :-  
    M12 is M1 - 1,  
    K12 is K1 - 1,  
    Mr2 is Mr + 1,  
    Kr2 is Kr + 1,  
    zust(M12,K12,'R',Mr2,Kr2).
```

```
% 1 Missionar von Links nach Rechts
zug(n(M1,K1,'L',Mr,Kr), n(M12,K1,'R',Mr2,Kr)) :-
    M12 is M1 - 1,
    Mr2 is Mr + 1,
    zust(M12,K1,'R',Mr2,Kr).
```

```
% 2 Missionare von Links nach Rechts
zug(n(M1,K1,'L',Mr,Kr), n(M12,K1,'R',Mr2,Kr)) :-
    M12 is M1 - 2,
    Mr2 is Mr + 2,
    zust(M12,K1,'R',Mr2,Kr).
```

Dann definieren wir die Überfahrten vom rechten zum linken Ufer:

```
% 1 Kannibale von Rechts nach Links
zug(n(M1,K1,'R',Mr,Kr), n(M1,K12,'L',Mr,Kr2)) :-
    K12 is K1 + 1,
    Kr2 is Kr - 1,
    zust(M1,K12,'L',Mr,Kr2).
```

```
% 2 Kannibalen von Rechts nach Links
zug(n(M1,K1,'R',Mr,Kr), n(M1,K12,'L',Mr,Kr2)) :-
    K12 is K1 + 2,
    Kr2 is Kr - 2,
    zust(M1,K12,'L',Mr,Kr2).
```

```
% 1 Kannibale, 1 Missionar von Rechts nach Links  
zug(n(M1,K1,'R',Mr,Kr), n(M12,K12,'L',Mr2,Kr2)) :-  
    M12 is M1 + 1,  
    K12 is K1 + 1,  
    Mr2 is Mr - 1,  
    Kr2 is Kr - 1,  
    zust(M12,K12,'L',Mr2,Kr2).
```

```
% 1 Missionar von Rechts nach Links
zug(n(M1,K1,'R',Mr,Kr), n(M12,K1,'L',Mr2,Kr)) :-
    M12 is M1 + 1,
    Mr2 is Mr - 1,
    zust(M12,K1,'L',Mr2,Kr).

% 2 Missionare von Rechts nach Links
zug(n(M1,K1,'R',Mr,Kr), n(M12,K1,'L',Mr2,Kr)) :-
    M12 is M1 + 2,
    Mr2 is Mr - 2,
    zust(M12,K1,'L',Mr2,Kr).
```

Nun definieren wir eine Zug-Folge, einen Pfad. Im einfachsten Fall erkennen wir keine doppelten Zustände, sondern beschränken nur die Suchtiefe, um „Endlosschleifen“ zu verhindern.

```
% Abbruch, wenn die Suchtiefe überschritten wurde:  
pfad(n(_,_,_,_), n(_,_,_,_), Tiefe, _) :-  
    Tiefe > 10, !, fail.
```

Um später einen Pfad ausgeben zu können, merken wir uns die gemachten Züge in einer Liste.

```
% Ein einzelner Zug ist ein zulässiger Pfad:  
pfad(n(M1,K1,P1,M2,K2), n(M3,K3,P3,M4,K4),  
    Tiefe, [n(M3,K3,P3,M4,K4)|[]]) :-  
    zug(n(M1,K1,P1,M2,K2), n(M3,K3,P3,M4,K4)).
```

Ein Pfad von A nach B existiert, falls es einen Zug von A nach C gibt und ein Pfad von C nach B existiert. Dabei müssen wir die Tiefe hochzählen:

```
pfad(n(M1,K1,P1,M2,K2), n(M5,K5,P5,M6,K6),  
      Tiefe, [n(M3,K3,P3,M4,K4)|L]) :-  
    T is Tiefe + 1,  
    zug(n(M1,K1,P1,M2,K2), n(M3,K3,P3,M4,K4)),  
    pfad(n(M3,K3,P3,M4,K4), n(M5,K5,P5,M6,K6), T, L).
```

Jetzt müssen wir das Ganze nur noch mit einem einfachen Aufruf versehen:

```
solve :- pfad(n(3,3,'L',0,0), n(0,0,'R',3,3), 0, L),  
          write([n(3,3,'L',0,0)|L])).
```

Können wir auf den Tiefenzähler verzichten?

- Wir werden die bereits untersuchten Zustände, also die bisherige Zugfolge, in einer Liste speichern.
- Vor der Ausführung eines neuen Zugs wird getestet, ob der Knoten in der Liste enthalten ist, also bereits untersucht wurde.
- Auf diese Weise vermeiden wir doppelte Zustände und verhindern „Endlosschleifen“.

Die Liste der bisherigen Züge wird in die Ergebnisliste kopiert, wenn eine Belegung erfolgreich war. Wir führen also zwei Listen:

- Die erste Liste enthält die Zustände, die bereits untersucht wurden. Wir füllen diese Liste beim rekursiven Abstieg.
- Da diese Informationen beim Rekursionsabbau wieder verloren gehen, kopieren wir die Werte am Rekursionsende in eine zweite Liste.

Am Rekursionsende wird die erste Liste in die zweite kopiert:

```
% Abbruch, wenn beide Zustände gleich sind:
```

```
pfad(n(M1,K1,P,M2,K2), n(M1,K1,P,M2,K2), L, L) :- !.
```

Sonst erzeugen wir einen Zug zu einem Zwischenziel,

- testen, ob der neue Zustand (das Zwischenziel) noch nicht in der ersten Liste gespeichert ist,
- hängen den neuen Zustand an die alte Liste an und
- versuchen, einen Pfad vom Zwischenziel zum Ziel zu finden.

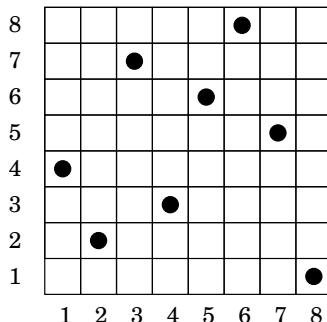
```
pfad(n(M1,K1,P1,M2,K2), n(M5,K5,P5,M6,K6), L1, L2) :-  
    zug(n(M1,K1,P1,M2,K2), n(M3,K3,P3,M4,K4)),  
    \+ member(n(M3,K3,P3,M4,K4), L1),  
    append(L1, [n(M3,K3,P3,M4,K4)], L),  
    pfad(n(M3,K3,P3,M4,K4), n(M5,K5,P5,M6,K6), L, L2).
```

- Allgemeines
- Erste Schritte
- Arithmetik
- Listen
- Backtracking
- Sortieren
- Input/Output
- Zahlenrätsel
- Kannibalen und Missionare
- *8-Damen-Problem*
- Schiebepuzzle
- A*-Algorithmus

8-Damen-Problem

Eine Lösung des Problems soll mittels des unären Prädikats `solution(Pos)` dargestellt werden.

Wir stellen den Zustand als Liste von x/y-Koordinaten dar:
`[X1/Y1, X2/Y2, X3/Y3, ..., X8/Y8]`.



Damit keine zwei Damen in der selben Spalte stehen, definieren wir ein Template `template([1/Y1, 2/Y2, 3/Y3, ..., 8/Y8])`.

Sammeln wir einige Eigenschaften von gültigen Lösungen:

- Eine leere Liste von Damen ist eine gültige Lösung, da sich keine Damen schlagen.

→ `solution([]).`

- Für eine nicht-leere Liste der Form `[X/Y, Others]` muss gelten:
 - Die Liste `Others` muss eine gültige Lösung sein.
 - `X` und `Y` müssen Integer-Werte zwischen 1 und 8 sein.
 - Die Dame auf Position `X/Y` darf sich mit keiner Dame der Liste `Others` schlagen.

→ `solution([X/Y | Others]) :-
 solution(Others),
 member(Y, [1,2,3,4,5,6,7,8]),
 noattack(X/Y, Others).`

8-Damen-Problem

Wie testen wir, ob sich die Dame auf Position X/Y mit denen der Liste Others schlägt, also ob `noattack(Q, QRest)` erfüllt ist?

- Wenn die Liste QRest leer ist, schlagen sich keine Damen.

→ `noattack(_, []).`

- Falls QRest die Form `[Q1 | QRest1]` hat,
 - darf die Dame Q die Dame Q1 nicht schlagen, und
 - die Dame Q darf die Damen in QRest1 nicht schlagen.

→ `noattack(X/Y, [X1/Y1 | Others]) :-`

```
    Y \= Y1,                % Zeilen ungleich
    Y1-Y \= X1-X,           % Diagonalen ungleich
    Y1-Y \= X-X1,
    noattack(X/Y, Others).
```

Als Abfrage erhalten wir also: `template(S), solution(S).`

- Allgemeines
- Erste Schritte
- Arithmetik
- Listen
- Backtracking
- Sortieren
- Input/Output
- Zahlenrätsel
- Kannibalen und Missionare
- 8-Damen-Problem
- *Schiebepuzzle*
- A*-Algorithmus

Zunächst nur eine naive Implementierung:

- Es werden keine Kreise erkannt.
- Vermeide endlose Rekursion durch Tiefenbeschränkung.
- Damit eine Lösung gefunden wird, müssen wir die Tiefe sukzessive erhöhen (iterative deepening).
- Die 0 beschreibt die Position der Lücke.

% Die beiden Endzustände:

```
puzzle(1,2,3,8,0,4,7,6,5,[],Tiefe).
```

```
puzzle(0,1,2,3,4,5,6,7,8,[],Tiefe).
```

% Abbruchbedingung für Tiefensuche:

```
puzzle(X1,X2,X3,Y1,Y2,Y3,Z1,Z2,Z3,L,Tiefe):-
```

```
    Tiefe > 7, !, fail.
```

```
% Zeile 1, Position links
% Lücke nach rechts schieben
puzzle(0,X2,X3,Y1,Y2,Y3,Z1,Z2,Z3,[X2|L],Tiefe) :-
    T is Tiefe+1,
    puzzle(X2,0,X3,Y1,Y2,Y3,Z1,Z2,Z3,L,T).

% Lücke nach unten schieben
puzzle(0,X2,X3,Y1,Y2,Y3,Z1,Z2,Z3,[Y1|L],Tiefe) :-
    T is Tiefe+1,
    puzzle(Y1,X2,X3,0,Y2,Y3,Z1,Z2,Z3,L,T).
```

```
% Zeile 1, Position mitte
% Lücke nach links schieben
puzzle(X1,0,X3,Y1,Y2,Y3,Z1,Z2,Z3,[X1|L],Tiefe) :-
    T is Tiefe+1,
    puzzle(0,X1,X3,Y1,Y2,Y3,Z1,Z2,Z3,L,T).

% Lücke nach rechts schieben
puzzle(X1,0,X3,Y1,Y2,Y3,Z1,Z2,Z3,[X3|L],Tiefe) :-
    T is Tiefe+1,
    puzzle(X1,X3,0,Y1,Y2,Y3,Z1,Z2,Z3,L,T).

% Lücke nach unten schieben
puzzle(X1,0,X3,Y1,Y2,Y3,Z1,Z2,Z3,[Y2|L],Tiefe) :-
    T is Tiefe+1,
    puzzle(X1,Y2,X3,Y1,0,Y3,Z1,Z2,Z3,L,T).
```

```
% Zeile 1, Position rechts
% Lücke nach links schieben
puzzle(X1,X2,0,Y1,Y2,Y3,Z1,Z2,Z3,[X2|L],Tiefe) :-
    T is Tiefe+1,
    puzzle(X1,0,X2,Y1,Y2,Y3,Z1,Z2,Z3,L,T).

% Lücke nach unten schieben
puzzle(X1,X2,0,Y1,Y2,Y3,Z1,Z2,Z3,[Y3|L],Tiefe) :-
    T is Tiefe+1,
    puzzle(X1,X2,Y3,Y1,Y2,0,Z1,Z2,Z3,L,T).
```

```
% Zeile 2, Position links
% Lücke nach rechts schieben
puzzle(X1,X2,X3,0,Y2,Y3,Z1,Z2,Z3,[Y2|L],Tiefe) :-
    T is Tiefe+1,
    puzzle(X1,X2,X3,Y2,0,Y3,Z1,Z2,Z3,L,T).

% Lücke nach oben schieben
puzzle(X1,X2,X3,0,Y2,Y3,Z1,Z2,Z3,[X1|L],Tiefe) :-
    T is Tiefe+1,
    puzzle(0,X2,X3,X1,Y2,Y3,Z1,Z2,Z3,L,T).

% Lücke nach unten schieben
puzzle(X1,X2,X3,0,Y2,Y3,Z1,Z2,Z3,[Z1|L],Tiefe) :-
    T is Tiefe+1,
    puzzle(X1,X2,X3,Z1,Y2,Y3,0,Z2,Z3,L,T).
```

```
% Zeile 2, Position mitte
% Lücke nach links schieben
puzzle(X1,X2,X3,Y1,0,Y3,Z1,Z2,Z3,[Y1|L],Tiefe) :-
    T is Tiefe+1,
    puzzle(X1,X2,X3,0,Y1,Y3,Z1,Z2,Z3,L,T).

% Lücke nach rechts schieben
puzzle(X1,X2,X3,Y1,0,Y3,Z1,Z2,Z3,[Y3|L],Tiefe) :-
    T is Tiefe+1,
    puzzle(X1,X2,X3,Y1,Y3,0,Z1,Z2,Z3,L,T).
```

```
% Zeile 2, Position mitte (Fortsetzung)
% Lücke nach oben schieben
puzzle(X1,X2,X3,Y1,0,Y3,Z1,Z2,Z3,[X2|L],Tiefe) :-
    T is Tiefe+1,
    puzzle(X1,0,X3,Y1,X2,Y3,Z1,Z2,Z3,L,T).

% Lücke nach unten schieben
puzzle(X1,X2,X3,Y1,0,Y3,Z1,Z2,Z3,[Z2|L],Tiefe) :-
    T is Tiefe+1,
    puzzle(X1,X2,X3,Y1,Z2,Y3,Z1,0,Z3,L,T).
```

```
% Zeile 2, Position rechts
% Lücke nach links schieben
puzzle(X1,X2,X3,Y1,Y2,0,Z1,Z2,Z3,[Y2|L],Tiefe) :-
    T is Tiefe+1,
    puzzle(X1,X2,X3,Y1,0,Y2,Z1,Z2,Z3,L,T).

% Lücke nach oben schieben
puzzle(X1,X2,X3,Y1,Y2,0,Z1,Z2,Z3,[X3|L],Tiefe) :-
    T is Tiefe+1,
    puzzle(X1,X2,0,Y1,Y2,X3,Z1,Z2,Z3,L,T).

% Lücke nach unten schieben
puzzle(X1,X2,X3,Y1,Y2,0,Z1,Z2,Z3,[Z3|L],Tiefe) :-
    T is Tiefe+1,
    puzzle(X1,X2,X3,Y1,Y2,Z3,Z1,Z2,0,L,T).
```

```
% Zeile 3, Position links
% Lücke nach rechts schieben
puzzle(X1,X2,X3,Y1,Y2,Y3,0,Z2,Z3,[Z2|L],Tiefe) :-
    T is Tiefe+1,
    puzzle(X1,X2,X3,Y1,Y2,Y3,Z2,0,Z3,L,T).

% Lücke nach oben schieben
puzzle(X1,X2,X3,Y1,Y2,Y3,0,Z2,Z3,[Y1|L],Tiefe) :-
    T is Tiefe+1,
    puzzle(X1,X2,X3,0,Y2,Y3,Y1,Z2,Z3,L,T).
```

```
% Zeile 3, Position mitte
% Lücke nach links schieben
puzzle(X1,X2,X3,Y1,Y2,Y3,Z1,0,Z3,[Z1|L],Tiefe) :-
    T is Tiefe+1,
    puzzle(X1,X2,X3,Y1,Y2,Y3,0,Z1,Z3,L,T).

% Lücke nach rechts schieben
puzzle(X1,X2,X3,Y1,Y2,Y3,Z1,0,Z3,[Z3|L],Tiefe) :-
    T is Tiefe+1,
    puzzle(X1,X2,X3,Y1,Y2,Y3,Z1,Z3,0,L,T).

% Lücke nach oben schieben
puzzle(X1,X2,X3,Y1,Y2,Y3,Z1,0,Z3,[Y2|L],Tiefe) :-
    T is Tiefe+1,
    puzzle(X1,X2,X3,Y1,0,Y3,Z1,Y2,Z3,L,T).
```

```
% Zeile 3, Position rechts
% Lücke nach links schieben
puzzle(X1,X2,X3,Y1,Y2,Y3,Z1,Z2,0,[Z2|L],Tiefe) :-
    T is Tiefe+1,
    puzzle(X1,X2,X3,Y1,Y2,Y3,Z1,0,Z2,L,T).

% Lücke nach oben schieben
puzzle(X1,X2,X3,Y1,Y2,Y3,Z1,Z2,0,[Y3|L],Tiefe) :-
    T is Tiefe+1,
    puzzle(X1,X2,X3,Y1,Y2,0,Z1,Z2,Y3,L,T).
```

Auch hier können wir, wie beim Kannibalen-Missionare-Problem, die Tiefenbeschränkung ersetzen, indem wir den neuen Zustand daraufhin überprüfen, ob er in der Liste der bereits besuchten Zustände enthalten ist.

Allerdings nützt uns das in diesem Fall gar nichts, da bei der Tiefensuche evtl. sehr ungünstige, d.h. sehr lange Wege gefunden werden und die Suchtiefe die Stack-Größe von Prolog überschreitet:

```
Fatal Error: global stack overflow
```

Daher werden wir den A^* -Algorithmus implementieren, um kürzere, bessere Lösungen zu finden. Das hier vorgestellte Programm ist angelehnt an die Lösung unter:

www.csupomona.edu/~jrfisher/www/prolog_tutorial/

- Allgemeines
- Erste Schritte
- Arithmetik
- Listen
- Backtracking
- Sortieren
- Input/Output
- Zahlenrätsel
- Kannibalen und Missionare
- 8-Damen-Problem
- Schiebepuzzle
- *A*-Algorithmus*

Einen Knoten des Baums stellen wir als Struktur `node(State, Depth, F, [Ancestors])` dar:

- State: repräsentierter Zustand
- Depth: Tiefe des Knotens im Baum, also $g(n)$.
- F: Wert der Bewertungsfunktion $f(n) = g(n) + h(n)$.
- Ancestors: Liste der Vorgänger des Knotens.

Um eine Lösung zu berechnen, rufen wir das Haupt-Prädikat des Programms `solve(Start, Sol)` auf. Dabei beschreibt Start den Startzustand und in Sol wird die Lösung dargestellt:

- Zunächst wird der f -Wert für Start berechnet.
- Anschließend wird mittels `search` ein Weg zum Ziel gesucht. Dabei wird die aktualisierte OpenList als Parameter übergeben.
- Schließlich muss die Vorgängerliste umgedreht werden.

```
solve(State,Solution) :-  
    f_function(State,0,F),  
    search([node(State,0,F,[])],S),  
    reverse(S,Solution).
```

Das Prädikat `f_function(State,D,F)` berechnet den Wert `F` für den Zustand `State` in Tiefe `D`.

- Dazu wird ein Prädikat `h_function(State,H)` aufgerufen, das wir für jedes Problem spezifisch implementieren müssen.
- Das Prädikat liefert in `H` den Wert der Schätzfunktion h vom Zustand `State` aus.

```
f_function(State,D,F) :-  
    h_function(State,H),  
    F is D + H.
```

Die Suche eines Zielknotens mittels `search` benötigt die OpenList:

- Der erste Knoten B wird aus der OpenList entfernt, alle Folgeknoten Children von B werden mittels `expand` bestimmt und mittels `insert_all` in die OpenList eingefügt. Anschließend wird die Suche rekursiv mit der aktualisierten OpenList aufgerufen.
- Die Suche endet, falls ein Ziel gefunden wird. Dazu muss für jedes Problem ein spezifisches Prädikat `goal(State)` bereit gestellt werden, das anzeigt, ob State ein gültiges Ziel ist.

```
search([node(State,_,_,Sol)|_], Sol) :- goal(State).  
search([B|Tail],S) :-  
    expand(B,Children),  
    insert_all(Children,Tail,Open),  
    search(Open,S).
```

Wenn ein Knoten `node(State, Depth, F, [A])` expandiert wird, passiert folgendes:

- Bestimme alle Kinder des Knotens.
- Jedes Kind hat Tiefe `Depth + 1`.
- Der F-Wert wird für jedes Kind berechnet.
- Die Vorgänger-Liste A wird vom Kind übernommen und ergänzt.

```
expand(node(State,D,_,S),All_My_Children) :-  
    bagof(node(Child,D1,F,[Move|S]),  
        ( D1 is D+1,  
          move(State,Child,Move),  
          f_function(Child,D1,F)),  
        All_My_Children).
```

Das Einfügen von Knoten in die OpenList erfolgt mittels einer Insertionsort-Variante. Dabei werden doppelte Zustände entfernt, die Einordnung der Elemente erfolgt anhand des f -Wertes.

```
insert_all([F|R],Open1,Open3) :-  
    insert(F,Open1,Open2),  
    insert_all(R,Open2,Open3).  
insert_all([],Open,Open).
```

```
insert(B,Open,Open) :- repeat_node(B,Open), ! .  
insert(B,[C|R],[B,C|R]) :- cheaper(B,C), ! .  
insert(B,[B1|R],[B1|S]) :- insert(B,R,S), ! .  
insert(B,[],[B]).
```

```
repeat_node(node(P,_,_,_), [node(P,_,_,_)|_]).  
cheaper(node(_,_,F1,_), node(_,_,F2,_)) :- F1 < F2.
```

8-Puzzle: Zunächst die Aufzählung aller möglichen Züge.

```
left( A/O/C/D/E/F/H/I/J , O/A/C/D/E/F/H/I/J ).  
left( A/B/C/D/O/F/H/I/J , A/B/C/O/D/F/H/I/J ).  
left( A/B/C/D/E/F/H/O/J , A/B/C/D/E/F/O/H/J ).  
left( A/B/O/D/E/F/H/I/J , A/O/B/D/E/F/H/I/J ).  
left( A/B/C/D/E/O/H/I/J , A/B/C/D/O/E/H/I/J ).  
left( A/B/C/D/E/F/H/I/O , A/B/C/D/E/F/H/O/I ).
```

```
up( A/B/C/O/E/F/H/I/J , O/B/C/A/E/F/H/I/J ).  
up( A/B/C/D/O/F/H/I/J , A/O/C/D/B/F/H/I/J ).  
up( A/B/C/D/E/O/H/I/J , A/B/O/D/E/C/H/I/J ).  
up( A/B/C/D/E/F/O/I/J , A/B/C/O/E/F/D/I/J ).  
up( A/B/C/D/E/F/H/O/J , A/B/C/D/O/F/H/E/J ).  
up( A/B/C/D/E/F/H/I/O , A/B/C/D/E/O/H/I/F ).
```

```
right( A/O/C/D/E/F/H/I/J , A/C/O/D/E/F/H/I/J ).  
right( A/B/C/D/O/F/H/I/J , A/B/C/D/F/O/H/I/J ).  
right( A/B/C/D/E/F/H/O/J , A/B/C/D/E/F/H/J/O ).  
right( O/B/C/D/E/F/H/I/J , B/O/C/D/E/F/H/I/J ).  
right( A/B/C/O/E/F/H/I/J , A/B/C/E/O/F/H/I/J ).  
right( A/B/C/D/E/F/O/I/J , A/B/C/D/E/F/I/O/J ).
```

```
down( A/B/C/O/E/F/H/I/J , A/B/C/H/E/F/O/I/J ).  
down( A/B/C/D/O/F/H/I/J , A/B/C/D/I/F/H/O/J ).  
down( A/B/C/D/E/O/H/I/J , A/B/C/D/E/J/H/I/O ).  
down( O/B/C/D/E/F/H/I/J , D/B/C/O/E/F/H/I/J ).  
down( A/O/C/D/E/F/H/I/J , A/E/C/D/O/F/H/I/J ).  
down( A/B/O/D/E/F/H/I/J , A/B/F/D/E/O/H/I/J ).
```

Nun müssen wir einige, für das 8-Puzzle spezifische Prädikate definieren. Als Schätzer verwenden wir die Manhattan-Distanz:

```
h_function(Puzz,H) :- p_fcn(Puzz,H).
```

```
p_fcn(A/B/C/D/E/F/G/H/I, P) :-  
    a(A,Pa), b(B,Pb), c(C,Pc),  
    d(D,Pd), e(E,Pe), f(F,Pf),  
    g(G,Pg), h(H,Ph), i(I,Pi),  
    P is Pa+Pb+Pc+Pd+Pe+Pf+Pg+Ph+Pg+Pi.
```

```
a(0,0). a(1,0). a(2,1). a(3,2). a(4,3).  
        a(5,4). a(6,3). a(7,2). a(8,1).  
b(0,0). b(1,1). b(2,0). b(3,1). b(4,2).  
        b(5,3). b(6,2). b(7,3). b(8,2).
```

c(0,0) . c(1,2) . c(2,1) . c(3,0) . c(4,1) .
c(5,2) . c(6,3) . c(7,4) . c(8,3) .
d(0,0) . d(1,1) . d(2,2) . d(3,3) . d(4,2) .
d(5,3) . d(6,2) . d(7,2) . d(8,0) .
e(0,0) . e(1,2) . e(2,1) . e(3,2) . e(4,1) .
e(5,2) . e(6,1) . e(7,2) . e(8,1) .
f(0,0) . f(1,3) . f(2,2) . f(3,1) . f(4,0) .
f(5,1) . f(6,2) . f(7,3) . f(8,2) .
g(0,0) . g(1,2) . g(2,3) . g(3,4) . g(4,3) .
g(5,2) . g(6,2) . g(7,0) . g(8,1) .
h(0,0) . h(1,3) . h(2,3) . h(3,3) . h(4,2) .
h(5,1) . h(6,0) . h(7,1) . h(8,2) .
i(0,0) . i(1,4) . i(2,3) . i(3,2) . i(4,1) .
i(5,0) . i(6,1) . i(7,2) . i(8,3) .

Das Prädikat `expand` des A*-Algorithmus benötigt das Prädikat `move`, das wir für das 8-Puzzle wie folgt definieren:

```
move(State, Child, l) :- left(State, Child).  
move(State, Child, u) :- up(State, Child).  
move(State, Child, r) :- right(State, Child).  
move(State, Child, d) :- down(State, Child).
```

In der Ergebnisliste, die dem Benutzer des Prolog-Programms angezeigt wird, stehen nur die Werte 'l', 'u', 'r' und 'd'.

Als Letztes müssen wir noch das Ziel-Prädikat definieren:

```
goal(1/2/3/8/0/4/7/6/5).
```

Jetzt können wir eine Lösung zu einem Startzustand berechnen lassen:

```
| ?- solve(1/6/2/7/0/3/5/8/4, S).  
S = [d,l,u,r,u,r,d,d,l,u] ? ;  
S = [d,l,u,r,u,r,d,d,l,u,l,r] ? ;  
S = [d,l,u,d,u,r,u,r,d,d,l,u] ? ;  
S = [d,l,u,r,u,r,d,d,l,u,l,r,l,r] ? ;  
S = [d,l,r,l,u,r,u,r,d,d,l,u] ?  
.....
```

Wir nutzen oben das Systemprädikat `bagof(X,P,L)`:

- Liefert die Liste L, die alle Objekte X enthält, die das Ziel P erfüllen.
- Macht in der Regel nur Sinn, falls X und P gemeinsame Variablen enthalten.
- Beispiel: Klassifizierung von Buchstaben in Vokale und Konsonanten.

```
class(a, vok).      class(d, kon).  
class(b, kon).      class(e, vok).  
class(c, kon).      class(f, kon).
```

Wir können mittels `bagof` eine Liste aller Vokale erstellen lassen:

```
| ?- bagof(Letter, class(Letter, vok), Letters).  
Letters = [a,e]  
yes
```

Weitere Systemprädikate:

- `integer(A)` wird einmal wahr, falls A eine ganze Zahl ist.
- `var(A)` wird einmal wahr, wenn A eine (nicht instantiierte) Variable ist.
- `repeat` wird beliebig oft wahr.
- und viele weitere

Anwendungen von Prolog: Kneipenführer für Darmstadt.

www.d120.de/kneipenfuehrer/kfv6/andere/KF12h.html