



هوش مصنوعی

UnderGrad.Q5/7: Artificial Intelligence

فريا نصیری مفخم

Faria Nassiri-Mofakham

4016282-01

<https://eng.ui.ac.ir/fnasiri>

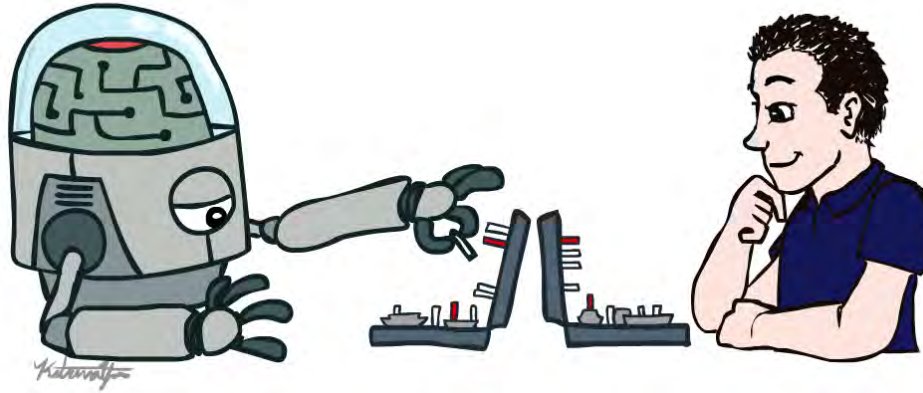
یکشنبه‌ها ۱۸-۱۷-۱۶، سه‌شنبه‌ها ۱۰-۸

lms.ui.ac.ir, ۳۶۰-۷

نیمسال ۴۰۴۲

CS 188: Artificial Intelligence

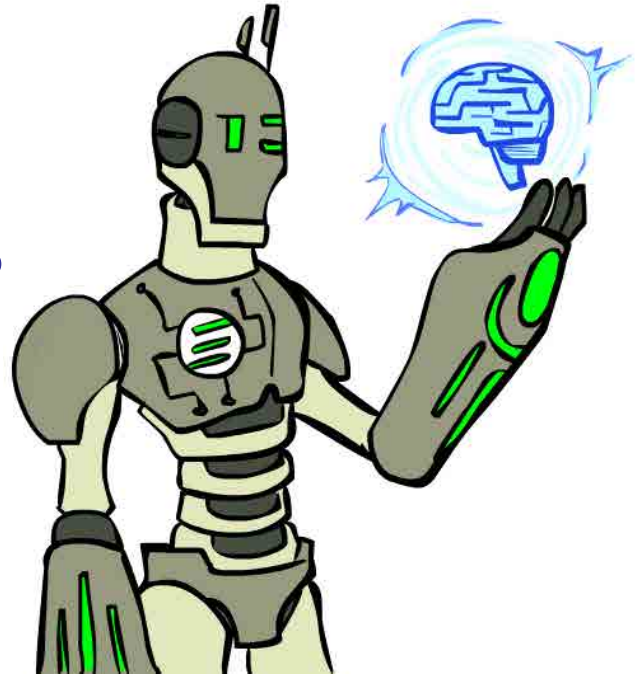
Introduction



Instructors: Stuart Russell and Dawn Song

Today

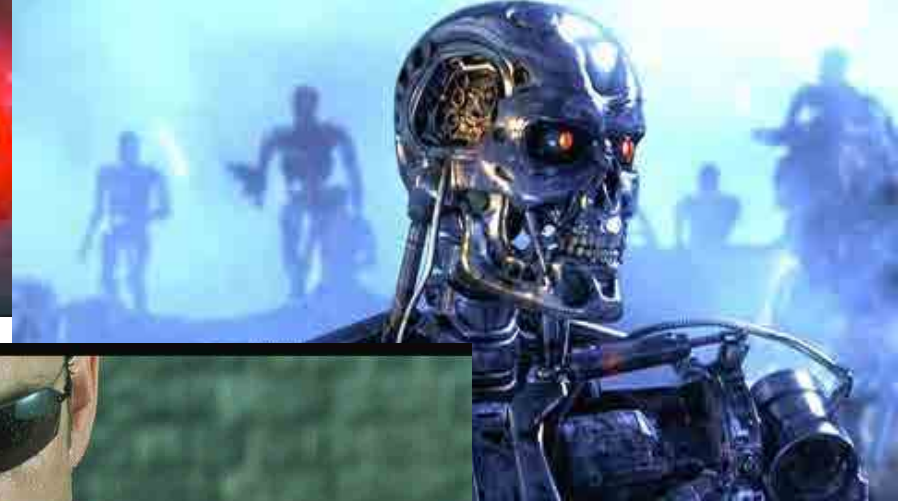
- What is artificial intelligence?
- Past: how did the ideas in AI come about?
- Present: what is the state of the art?
- Future: will robots take over the world?



Movie AI



Movie AI



YESTERDAY DR. WILL CASTER WAS ONLY HUMAN

JOHNNY DEPP REBECCA HALL PAUL BETTANY KATE MARA CILLIAN MURPHY AND MORGAN FREEMAN

TRANSCENDENCE

ALCON ENTERTAINMENT PRESENTS A WOLFGANG PETERSEN FILM "TRANSCENDENCE" STARRING JOHNNY DEPP, MORGAN FREEMAN, REBECCA HALL, KATE MARA, CILLIAN MURPHY, PAUL BETTANY, AND PAUL DOOLEY. MUSIC BY JAMES NEWTON HOWARD. COSTUME DESIGNER: JESSICA KAPLAN. EDITOR: JAMES HARRIS. EXECUTIVE PRODUCERS: JAMES HARRIS, JAMES HARRIS, JAMES HARRIS. PRODUCED BY JAMES HARRIS. WRITTEN BY JAMES HARRIS. DIRECTED BY WOLFGANG PETERSEN.

entertainmentfilms.co.uk

News AI

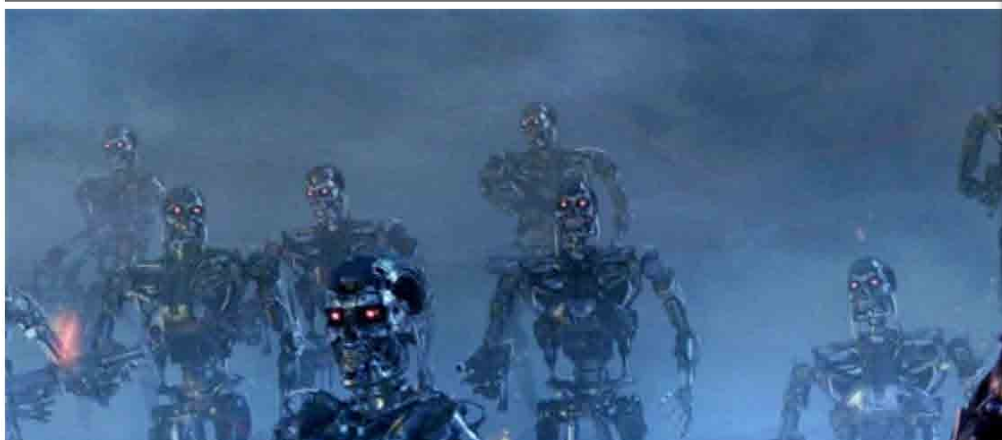
AI is the biggest risk we face as a civilisation, Elon Musk says

Billionaire burn: Musk says Zuckerberg's understanding of AI threat 'is limited'

HOME » FINANCE » FINANCE TOPICS » DAVOS

'Sociopathic' robots could overrun the human race in a generation

Computers should be trained to serve humans to reduce their threat to the human race, says a leading expert on artificial intelligence



LIVESCIENCE

NEWS TECH HEALTH PLANET EARTH

Live Science > Tech

Lifelike 'Sophia' Robot Granted Citizenship to Saudi Arabia

By Mindy Weisberger, Senior Writer | October 30, 2017 03:39pm ET



MORE ▾

Be the most interesting person you know, subscribe to LiveScience.

Subscribe >



News AI

TECH • ARTIFICIAL INTELLIGENCE

United Kingdom Plans \$1.3 Billion Intelligence Push

France to spend \$1.8 billion on compete with U.S., China

EU wants to invest £18b development

China's Got a Huge Artificial Intelligence Plan

'Whoever leads in AI will rule the world': Putin to Russian children on Knowledge Day

Published time: 1 Sep, 2017 14:08

Edited time: 1 Sep, 2017 14:40



News AI

NATURAL 'PROZAC': DOES IT REALLY WORK?

IBM's Watson Jeopardy Computer Shuts Down Humans in Final Game

DAILY NEWS 9 March 2016

Sili

'I'm in shock!' How world's best hum



Google DeepMind
Challenge Match
8 - 15 March 2016

Who is Stoker?

(I FOR ONE WELCOME OUR
NEW COMPUTER OVERLORDS)

Blizzard will show off Google's Deepmind AI in StarCraft 2 later this week

By Andy Chalk 4 hours ago

Google and Blizzard launched the artificial intelligence project in 2016.

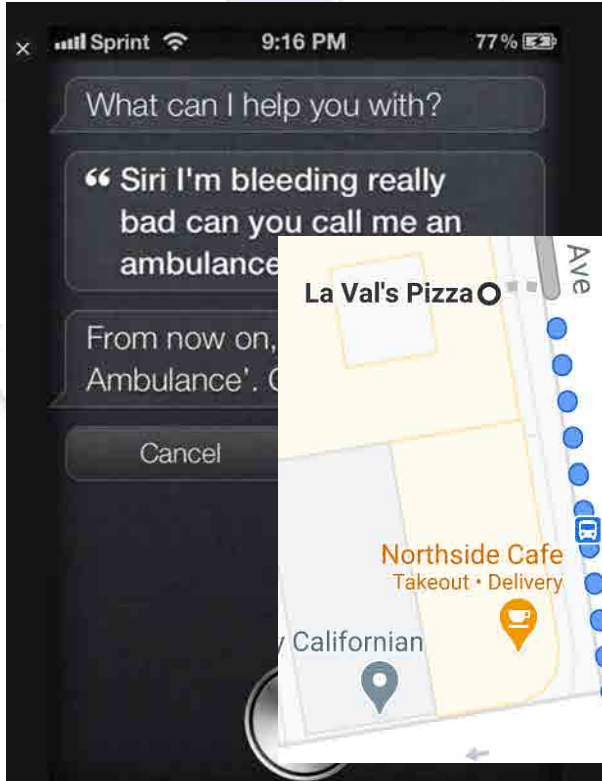


COMMENTS



Real AI

Google Translate



ENGLISH

SPANISH

CHINESE (TRADITIONAL)

ENGLISH

FRENCH



a Hall

136



TUG
CAUTION
MAY CONTAIN
CHEMOTHERAPY DRUG

CAUTION
MAY CONTAIN
CHEMOTHERAPY DRUG

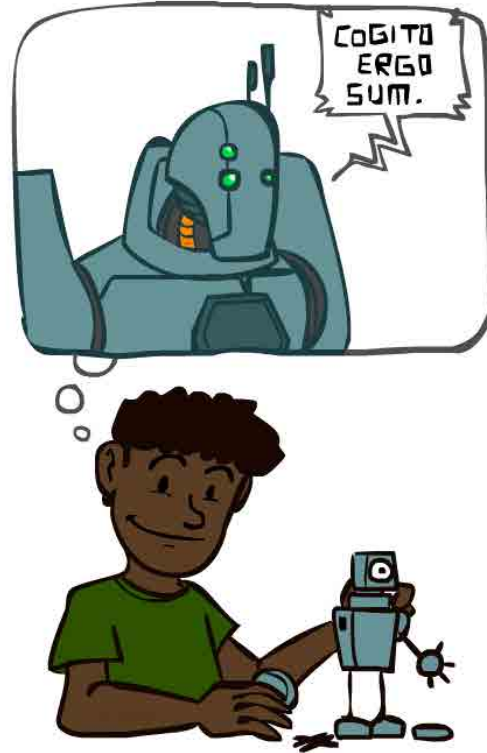




Boston Dynamics



A (Short) History of AI



A short prehistory of AI

■ Prehistory:

- **Philosophy** (reasoning, planning, learning, science, automation)

- Aristotle: For if every instrument could accomplish its own work, obeying or anticipating the will of others . . . if, in like manner, the shuttle would weave and the plectrum touch the lyre without a hand to guide them, chief workmen would not want servants, nor masters slaves

- **Psychology** (learning, cognitive models)

- **Linguistics** (grammars, formal representation of meaning)

■ Near miss (1842):

- Babbage design for universal machine
- Lovelace: “a thinking machine” for “all subjects in the universe.”

AI's official birth: Dartmouth, 1956

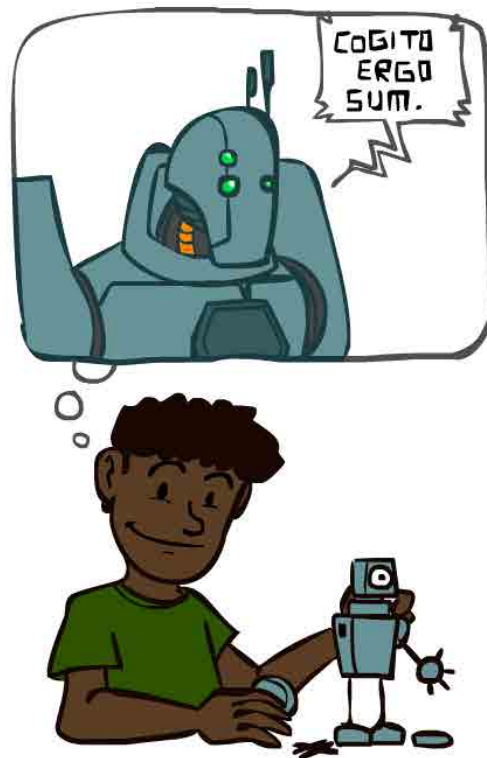


“An attempt will be made to find how to make machines use language, form abstractions and concepts, solve kinds of problems now reserved for humans, and improve themselves. ***We think that a significant advance can be made if we work on it together for a summer.***”

**John McCarthy and Claude Shannon
Dartmouth Workshop Proposal**

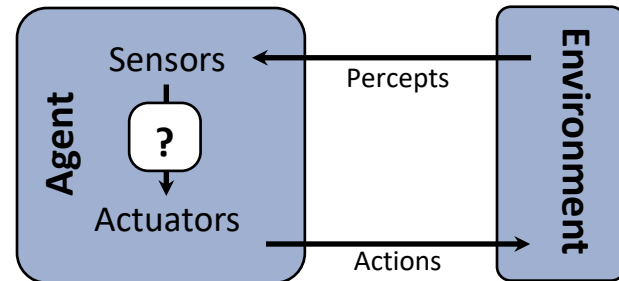
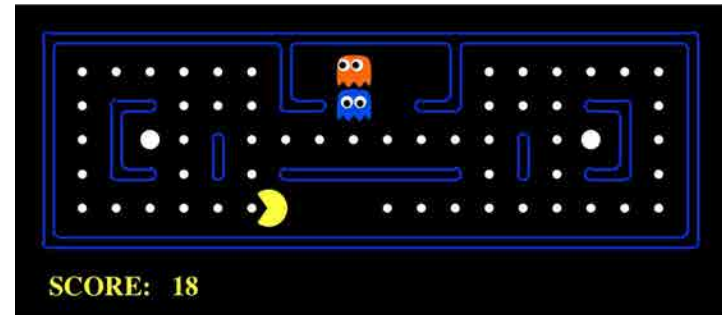
A (Short) History of AI

- 1940-1950: Early days
 - 1943: McCulloch & Pitts: Boolean circuit model of brain
 - 1950: Turing's "Computing Machinery and Intelligence"
- 1950—70: Excitement: Look, Ma, no hands!
 - 1950s: Early AI programs: chess, checkers (RL), theorem proving
 - 1956: Dartmouth meeting: "Artificial Intelligence" adopted
 - 1965: Robinson's complete algorithm for logical reasoning
- 1970—90: Knowledge-based approaches
 - 1969—79: Early development of knowledge-based systems
 - 1980—88: Expert systems industry booms
 - 1988—93: Expert systems industry busts: "AI Winter"
- 1990— 2012: Statistical approaches + subfield expertise
 - Resurgence of probability, focus on uncertainty
 - General increase in technical depth
 - Agents and learning systems... "AI Spring"?
- 2012— ____: Excitement: Look, Ma, no hands again?
 - Big data, big compute, deep learning
 - AI used in many industries



AI as Designing Rational Agents

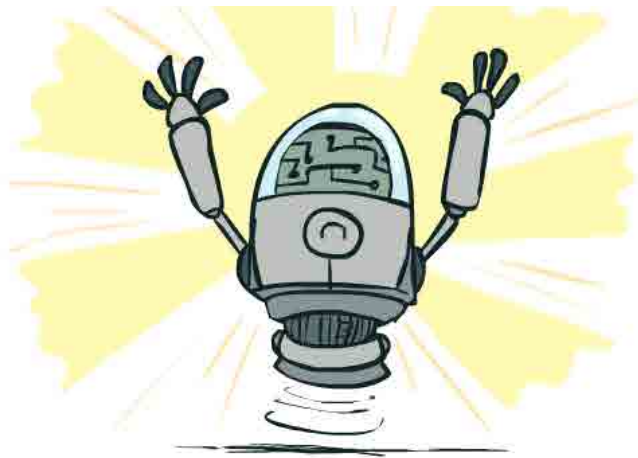
- An **agent** is an entity that *perceives* and *acts*.
- A **rational agent** selects actions that maximize its expected utility.
- Characteristics of the **sensors, actuators, and environment** dictate techniques for selecting rational actions
- **This course** is about:
 - General AI techniques for many problem types
 - Learning to choose and apply the technique appropriate for each problem



What Can AI Do?

Quiz: Which of the following can be done at present?

- ✓ Play a decent game of table tennis?
- ✓ Play a decent game of Jeopardy?
- ✓ Drive safely along a curving mountain road?
- ✗ Drive safely along Telegraph Avenue?
- ✓ Buy a week's worth of groceries on the web?
- ✗ Buy a week's worth of groceries at Berkeley Bowl?
- ? Discover and prove a new mathematical theorem?
- ✗ Converse successfully with another person for an hour?
- ? Perform a surgical operation?
- ✓ Translate spoken Chinese into spoken English in real time?
- ? Fold the laundry and put away the dishes?
- ✗ Write an intentionally funny story?



Unintentionally Funny Stories

Once upon a time
and a vain crow
in his tree, hold
mouth. He not
piece of cheese
swallowed the
the crow. The



Janelle Shane

@JanelleCShane

Follow

Tried retraining the neural net on just "what do you get when you cross a X with a X?" jokes. Results did not improve. And for some reason, bungees are its favorite thing.

What do you get when you cross a dog and a vampire?

A bungee

What do you get when you cross a cow with a rhino?

A bungee with a dog

What do you get when you cross a street and a cow?

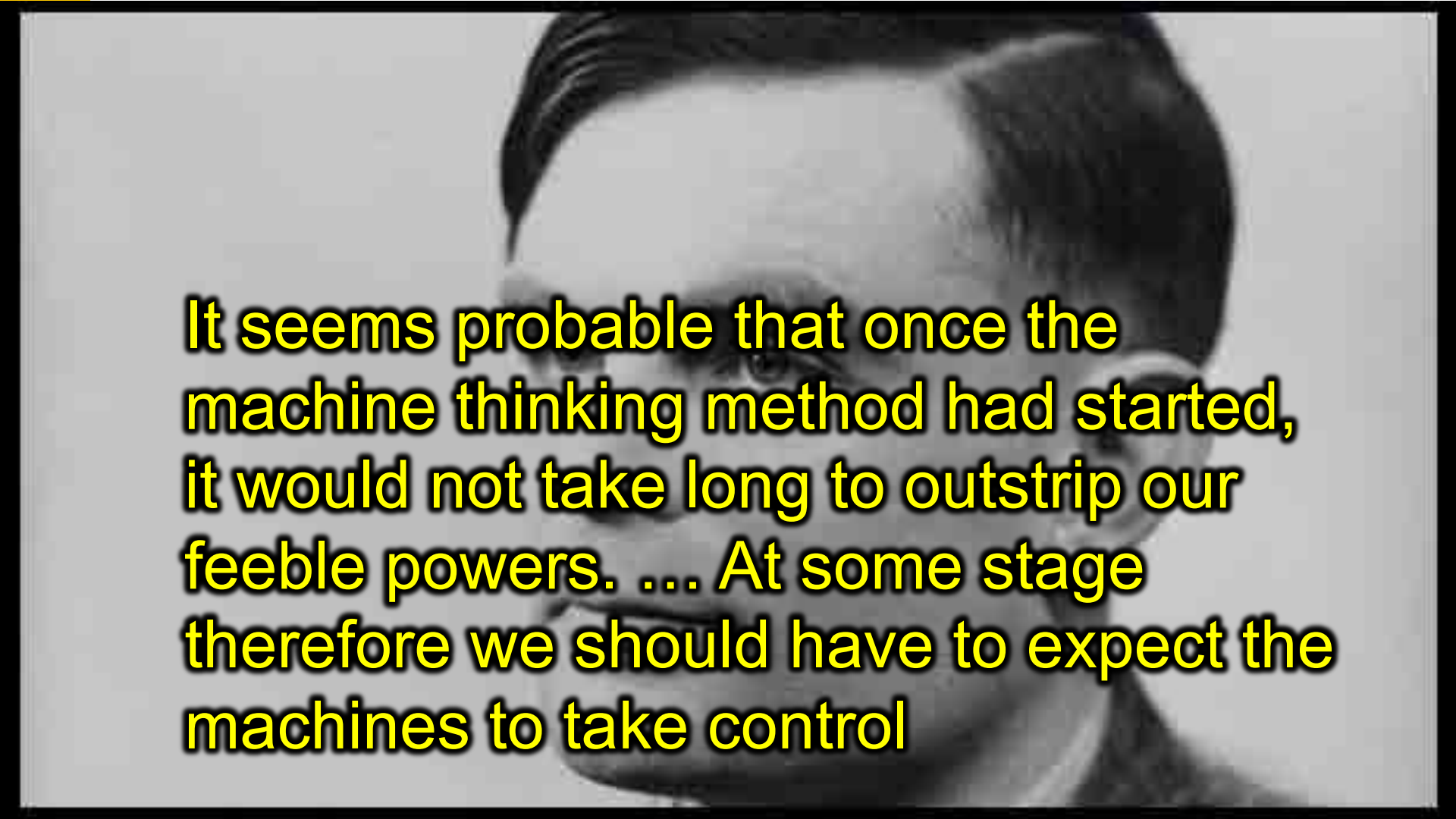
A bungee with a bungee and a rhino

What do you get when you cross a pig with a cow with a party?

Because the engineers with a dog

Future

- We are doing AI...
 - To create intelligent systems
 - The more intelligent, the better
 - To gain a better understanding of human intelligence
 - To magnify those benefits that flow from it
 - E.g., net present value of human-level AI $\geq$ \$13,500T
 - Might help us avoid war and ecological catastrophes, achieve immortality and expand throughout the universe
- What if we succeed?



It seems probable that once the machine thinking method had started, it would not take long to outstrip our feeble powers. ... At some stage therefore we should have to expect the machines to take control

What's bad about better AI?

- AI that is incredibly good at achieving something other than what we really want
- AI, economics, statistics, operations research, control theory all assume utility to be *fixed, known, and exogenously specified*
 - ~~Machines are intelligent to the extent that their actions can be expected to achieve their objectives~~
 - Machines are beneficial to the extent that their actions can be expected to achieve our objectives

What's bad about better AI?

- AI that is incredibly good at achieving something other than what we really want
- AI, economics, statistics, operations research, control theory all assume utility to be *fixed, known, and exogenously specified*
 - ~~Machines are intelligent to the extent that their actions can be expected to achieve their objectives~~
 - Machines are beneficial to the extent that their actions can be expected to achieve our objectives

A new model for AI

1. The machine's only objective is to maximize the realization of human preferences
2. The robot is initially uncertain about what those preferences are
3. Human behavior provides evidence about human preferences

“The essential task of our age” [Nick Bostrom, Professor of Philosophy, Oxford]

A new model for AI

1. The machine's only objective is to maximize the realization of human preferences
2. The robot is initially uncertain about what those preferences are
3. Human behavior provides evidence about human preferences

The standard model of AI is a special case, where the human can exactly and correctly program the objective into the machine

عامل هوشمند

	Humanly (مثل انسان)	Rationally (منطقی: ایده آل هوشمندی)	
Thinking humanly (فکر)	Cognitive Modeling Introspection, Psychological experiments, brain imaging	Laws of thought Logics, Probability	Thinking rationally
Acting Humanly (رفتار)	Turing Test NLP, Knowledge representation, Automated reasoning, Machine Learning	Rational agent Do the right thing	Acting rationally



1.2.1 Philosophy

- Can formal rules be used to draw valid conclusions?
- How does the mind arise from a physical brain?
- Where does knowledge come from?
- How does knowledge lead to action?

1.2.2 Mathematics

- What are the formal rules to draw valid conclusions?
- What can be computed?
- How do we reason with uncertain information?

1.2.3 Economics

- How should we make decisions in accordance with our preferences?
- How should we do this when others may not go along?
- How should we do this when the payoff may be far in the future?

1.2.4 Neuroscience

- How do brains process information?

1.2.5 Psychology

- How do humans and animals think and act?

1.2.6 Computer engineering

- How can we build an efficient computer?

1.2.7 Control theory and cybernetics

- How can artifacts operate under their own control?

1.2.8 Linguistics

- How does language relate to thought?

	Supercomputer	Personal Computer	Human Brain
Computational units	10 ⁶ GPUs + CPUs 10 ¹⁵ transistors	8 CPU cores 10 ¹⁰ transistors	10 ⁶ columns 10 ¹¹ neurons
Storage units	10 ¹⁶ bytes RAM 10 ¹⁷ bytes disk	10 ¹⁰ bytes RAM 10 ¹² bytes disk	10 ¹¹ neurons 10 ¹⁴ synapses
Cycle time	10 ⁻⁹ sec	10 ⁻⁹ sec	10 ⁻³ sec
Operations/sec	10 ¹⁸	10 ¹⁰	10 ¹⁷ 7

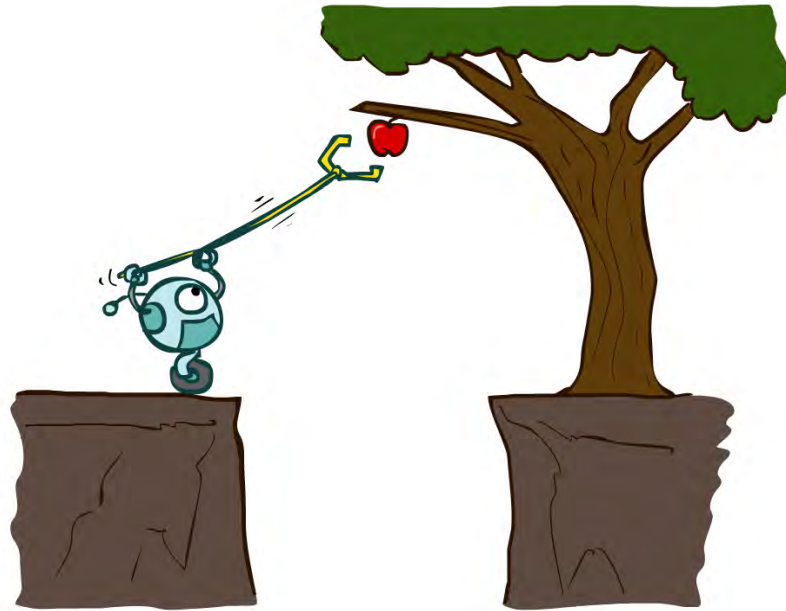
State of the art

- Robotic vehicles
- Legged locomotion
- Autonomous planning and scheduling
- Machine translation
- Speech recognition
- Recommendations
- Game playing
- Image understanding
- Medicine
- Climate science



CS 188: Artificial Intelligence

Agents and environments

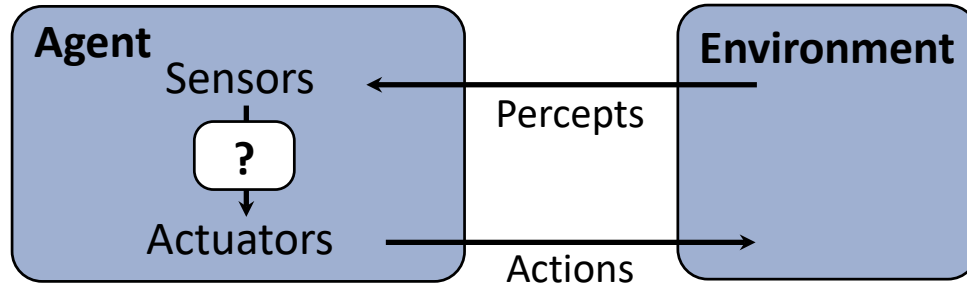


Instructors: Stuart Russell and Dawn Song

Outline

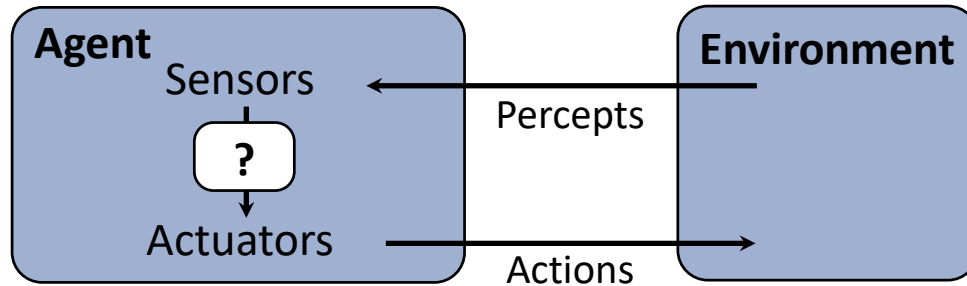
- Agents and environments
- Rationality
- PEAS (Performance measure, Environment, Actuators, Sensors)
- Environment types
- Agent types

Agents and environments



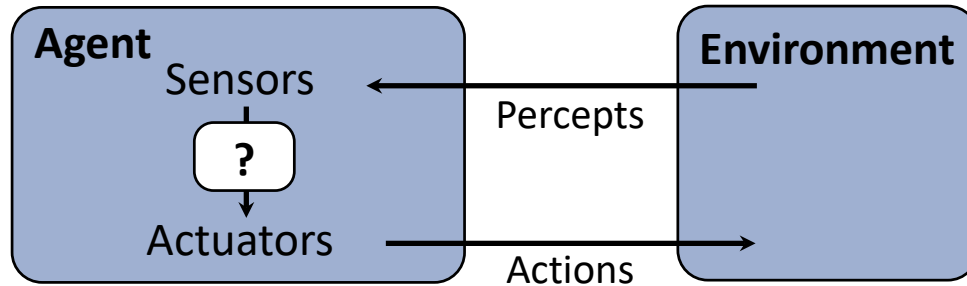
- An agent *perceives* its environment through *sensors* and *acts* upon it through *actuators* (or *effectors*, depending on whom you ask)

Agents and environments



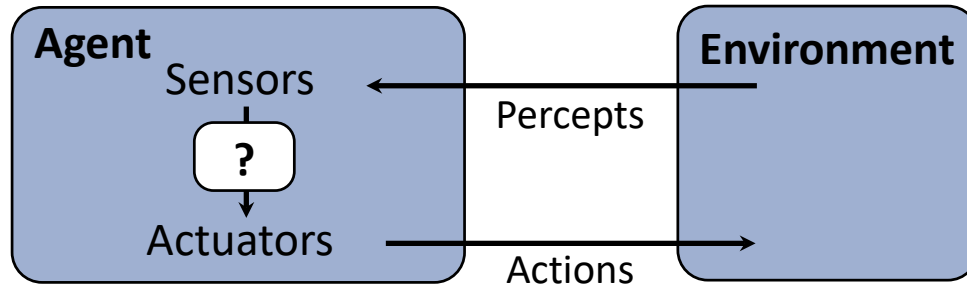
- Are humans agents?
- Yes!
 - Sensors = vision, audio, touch, smell, taste, proprioception
 - Actuators = muscles, secretions, changing brain state

Agents and environments



- Are pocket calculators agents?
- Yes!
 - Sensors = key state sensors
 - Actuators = digit display

Agents and environments



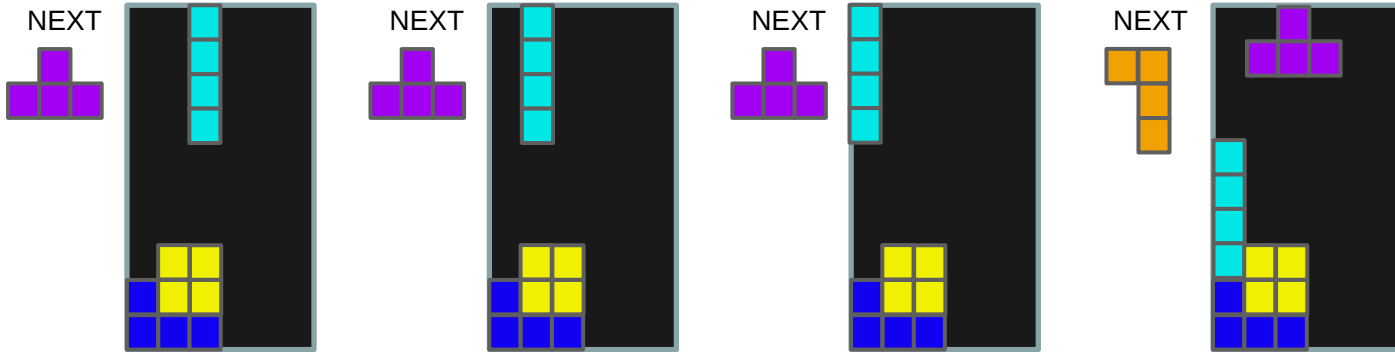
- AI is more interested in agents with large computational resources and environments that require nontrivial decision making



Agent functions

- The **agent function** maps from percept histories to actions:
 - $f: \mathcal{P}^* \rightarrow \mathcal{A}$
 - I.e., the agent's actual response to any sequence of percepts

Percept



Action

LEFT

LEFT

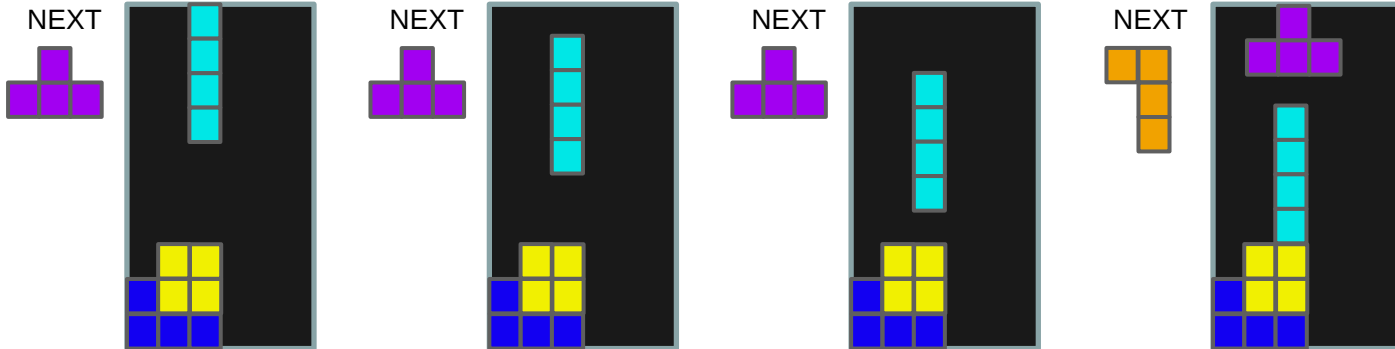
DROP

RIGHT

Agent programs

- The **agent program** I runs on some machine M to implement f :
 - $f = \text{Agent}(I, M)$
 - Real machines have limited speed and memory, introducing delay, so agent function f depends on M as well as I

Percept



NOOP

NOOP

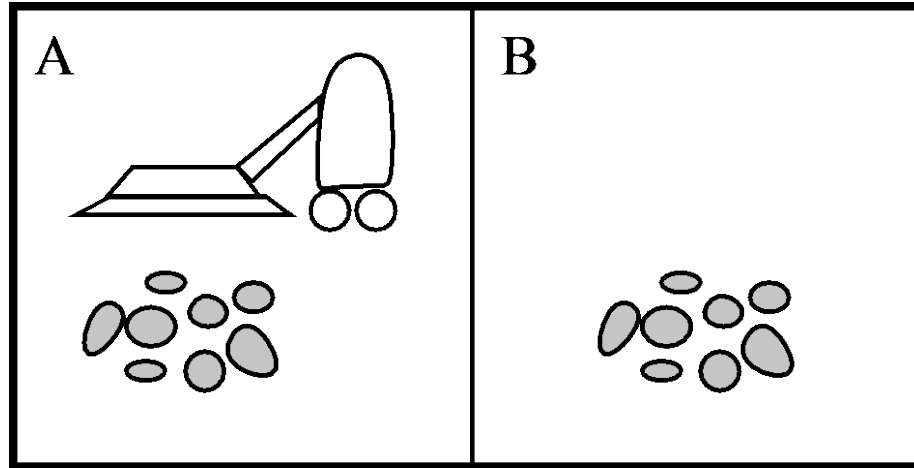
NOOP

LEFT

Agent functions and agent programs

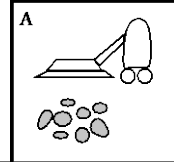
- Can every agent function be implemented by some agent program?
 - No! Consider agent for halting problems, NP-hard problems, chess with a slow PC

Example: Vacuum world



- Percepts: [location,status], e.g., [A,*Dirty*]
- Actions: *Left, Right, Suck, NoOp*

Vacuum cleaner agent



Agent function

Percept sequence	Action
[A,Clean]	Right
[A,Dirty]	Suck
[B,Clean]	Left
[B,Dirty]	Suck
[A,Clean],[B,Clean]	Left
[A,Clean],[B,Dirty]	Suck
etc	etc

Agent program

```
function Reflex-Vacuum-Agent([location,status])  
    returns an action  
if status = Dirty then return Suck  
else if location = A then return Right  
else if location = B then return Left
```

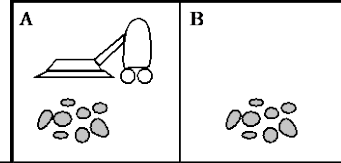
What is the *right* agent function?

Can it be implemented by a small agent program?

(Can we ask, “What is the right agent program?”)



Rationality



- Fixed **performance measure** evaluates the environment sequence
 - one point per square cleaned up?
 - NO! Rewards an agent who dumps dirt and cleans it up
 - one point per clean square per time step, for $t = 1, \dots, T$
- A **rational agent** chooses whichever action maximizes the **expected** value of the performance measure
 - given the percept sequence to date and prior knowledge of environment

Does Reflex-Vacuum-Agent implement a rational agent function?

Yes, if movement is free, or new dirt arrives frequently

Rationality, contd.

- Are rational agents *omniscient*?
 - No – they are limited by the available percepts
- Are rational agents *clairvoyant*?
 - No – they may lack knowledge of the environment dynamics
- Do rational agents *explore* and *learn*?
 - Yes – in unknown environments these are essential
- Do rational agents *make mistakes*?
 - No – but their actions may be unsuccessful
- Are rational agents *autonomous* (i.e., transcend initial program)?
 - Yes – as they learn, their behavior depends more on their own experience

The task environment - PEAS

- Performance measure

- -1 per step; + 10 food; +500 win; -500 die;
+200 hit scared ghost

- Environment

- Pacman dynamics (incl ghost behavior)

- Actuators

- Left Right Up Down

- Sensors

- Entire state is visible (except power pellet duration)



PEAS: Automated taxi

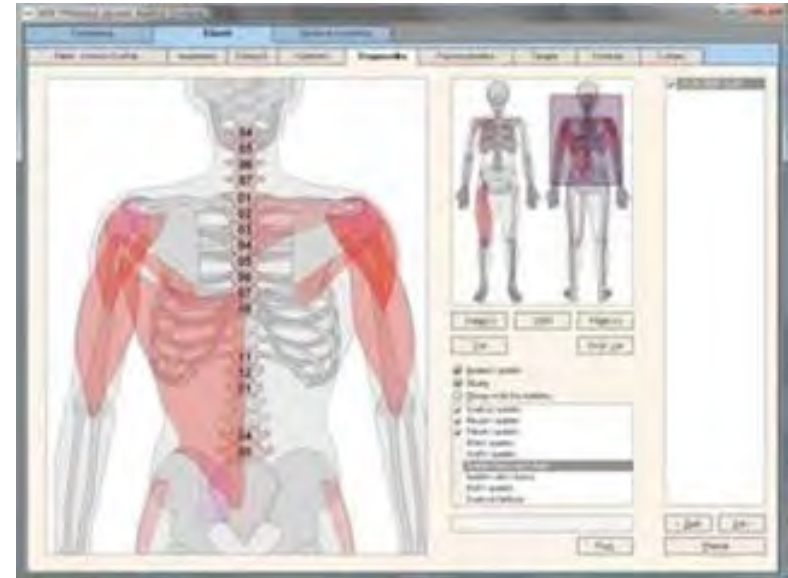
- Performance measure
 - Income, happy customer, vehicle costs, fines, insurance premiums
- Environment
 - US streets, other drivers, customers, weather, police...
- Actuators
 - Steering, brake, gas, display/speaker
- Sensors
 - Camera, radar, accelerometer, engine sensors, microphone, GPS



Image: <http://nypost.com/2014/06/21/how-google-might-put-taxi-drivers-out-of-business/>

PEAS: Medical diagnosis system

- Performance measure
 - Patient health, cost, reputation
- Environment
 - Patients, medical staff, insurers, courts
- Actuators
 - Screen display, email
- Sensors
 - Keyboard/mouse



Environment types

	Pacman	Backgammon	Diagnosis	Taxi
Fully or partially observable				
Single-agent or multiagent				
Deterministic or stochastic				
Static or dynamic				
Discrete or continuous				
Known physics?				
Known perf. measure?				

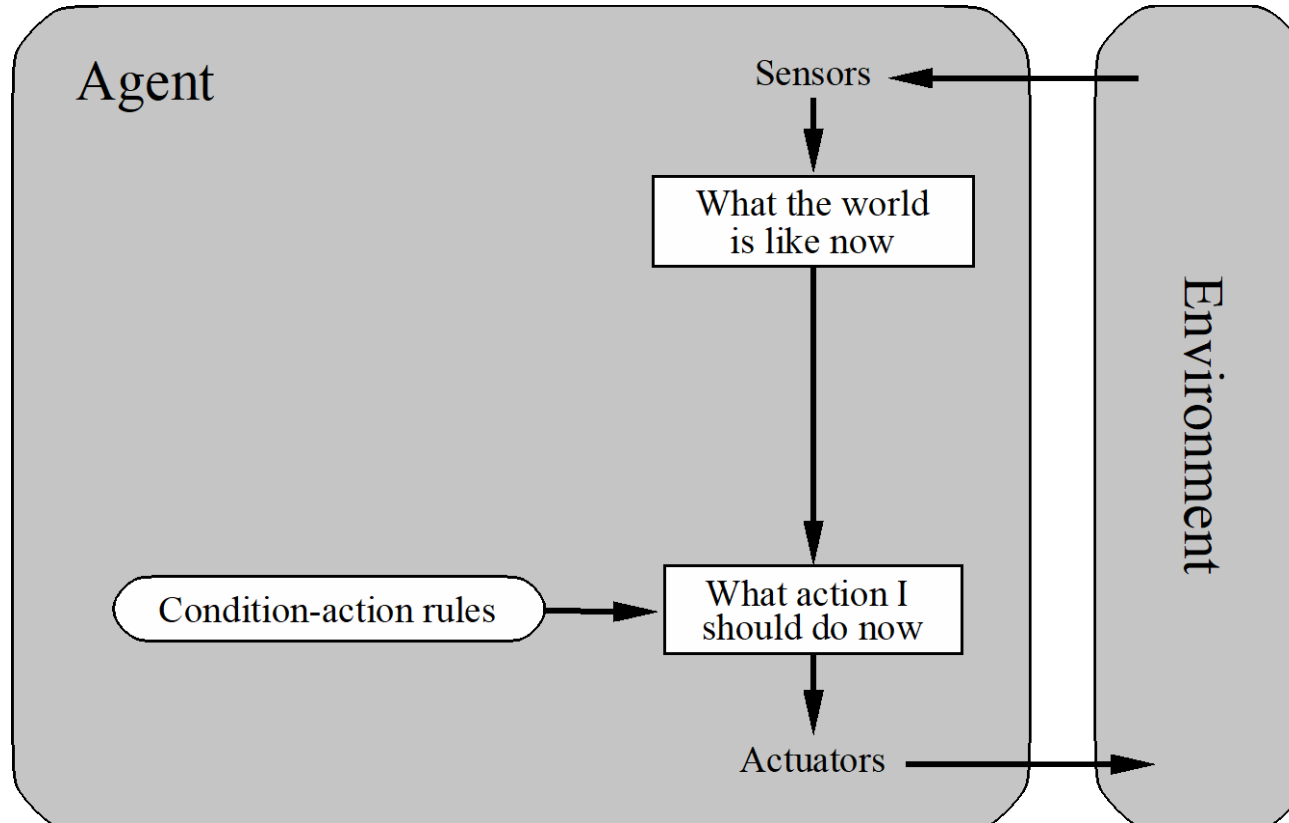
Agent design

- The environment type largely determines the agent design
 - *Partially observable* => agent requires *memory* (internal state)
 - *Stochastic* => agent may have to prepare for *contingencies*
 - *Multi-agent* => agent may need to behave *randomly*
 - *Static* => agent has time to compute a rational decision
 - *Continuous time* => continuously operating *controller*
 - *Unknown physics* => need for *exploration*
 - *Unknown perf. measure* => observe/interact with *human principal*

Agent types

- In order of increasing generality and complexity
 - Simple reflex agents
 - Reflex agents with state
 - Goal-based agents
 - Utility-based agents

Simple reflex agents



Pacman agent in Python

```
class GoWestAgent(Agent):
```

```
    def getAction(self, percept):
```

```
        if Directions.WEST in percept.getLegalPacmanActions():
```

```
            return Directions.WEST
```

```
        else:
```

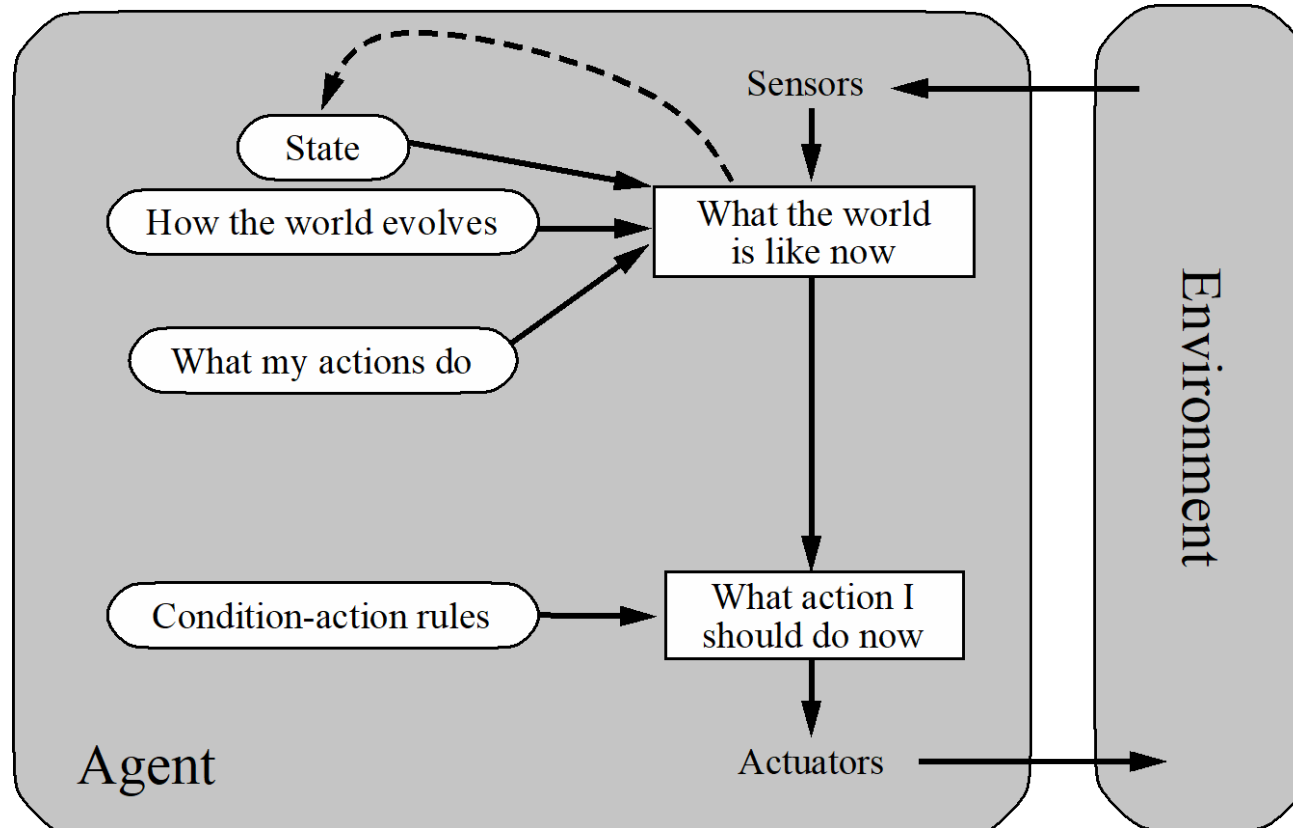
```
            return Directions.STOP
```



Pacman agent contd.

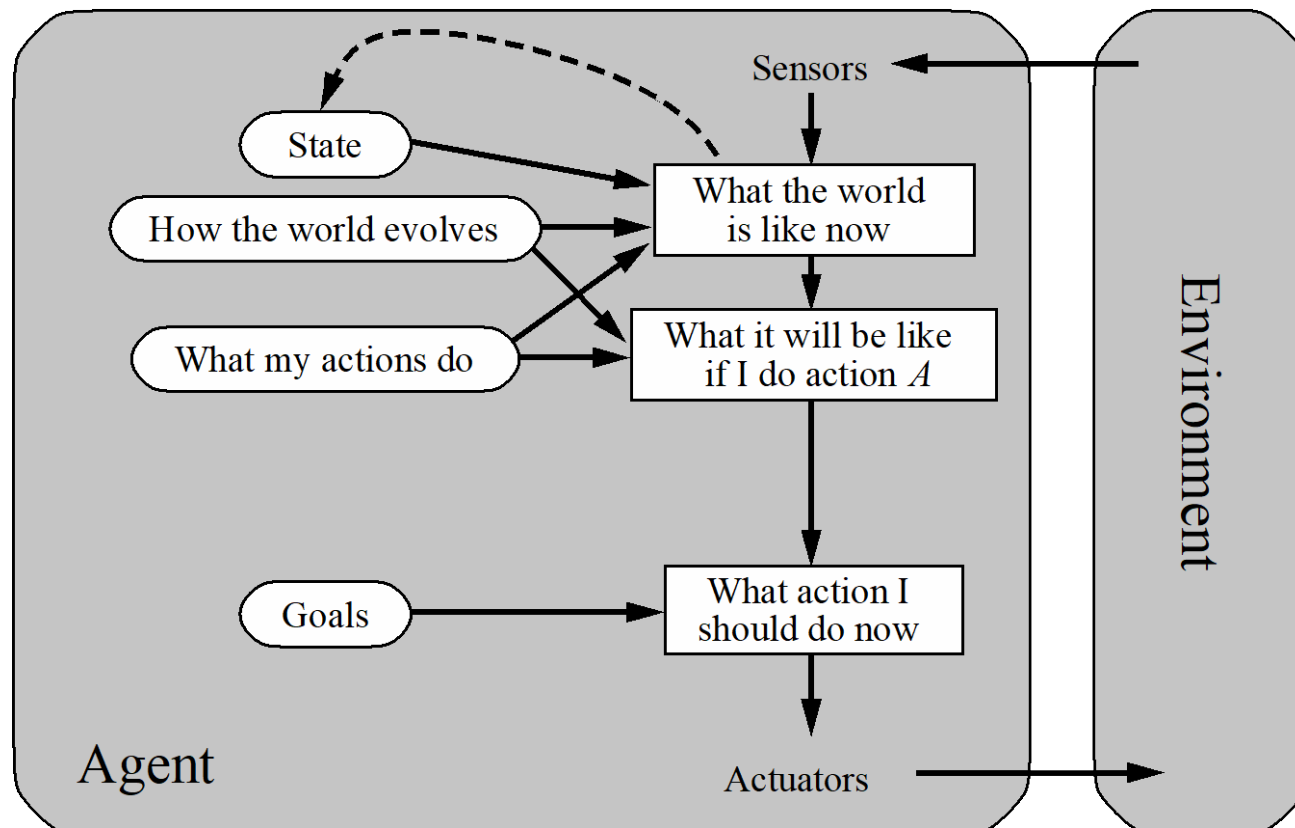
- Can we (in principle) extend this reflex agent to behave well in all standard Pacman environments?
 - No – Pacman is not quite fully observable (power pellet duration)
 - Otherwise, yes – we can (in principle) make a lookup table.....

Reflex agents with state



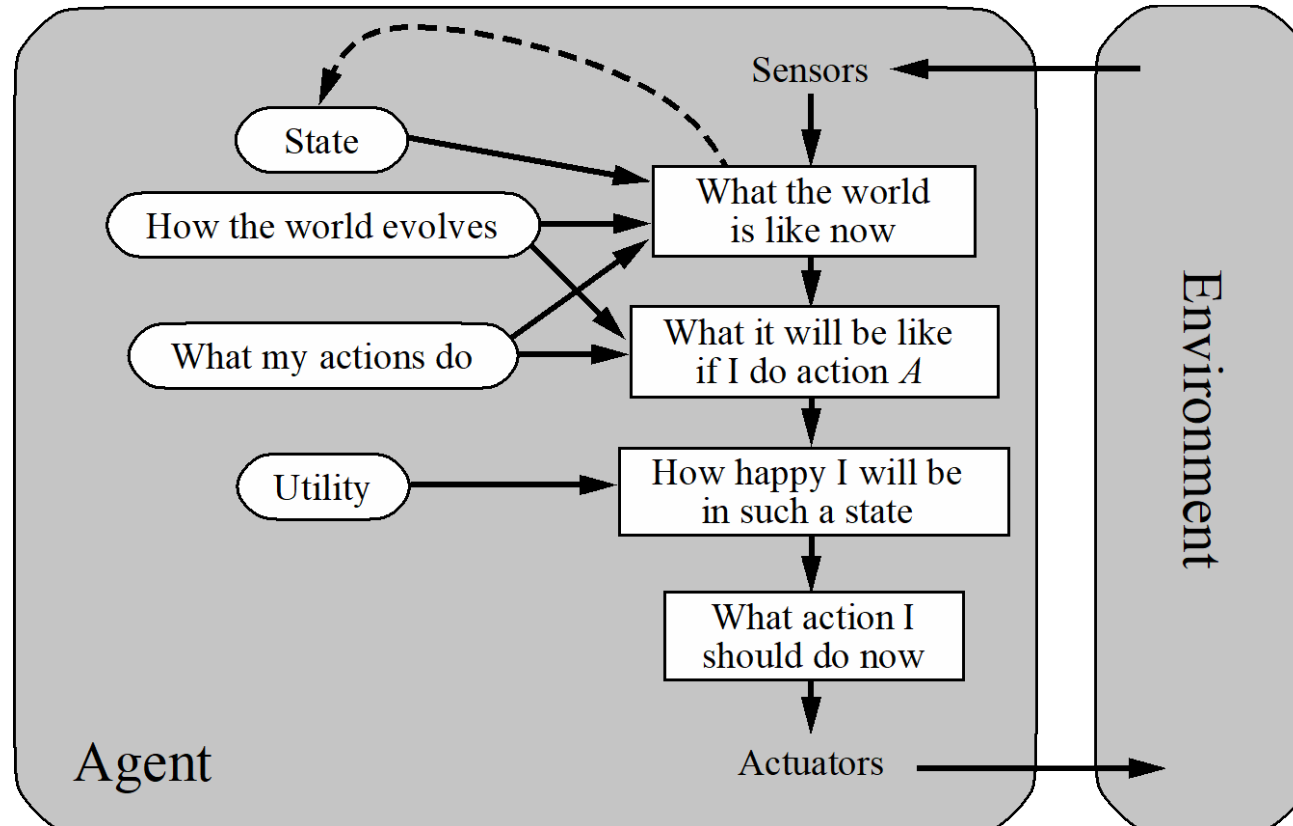


Goal-based agents





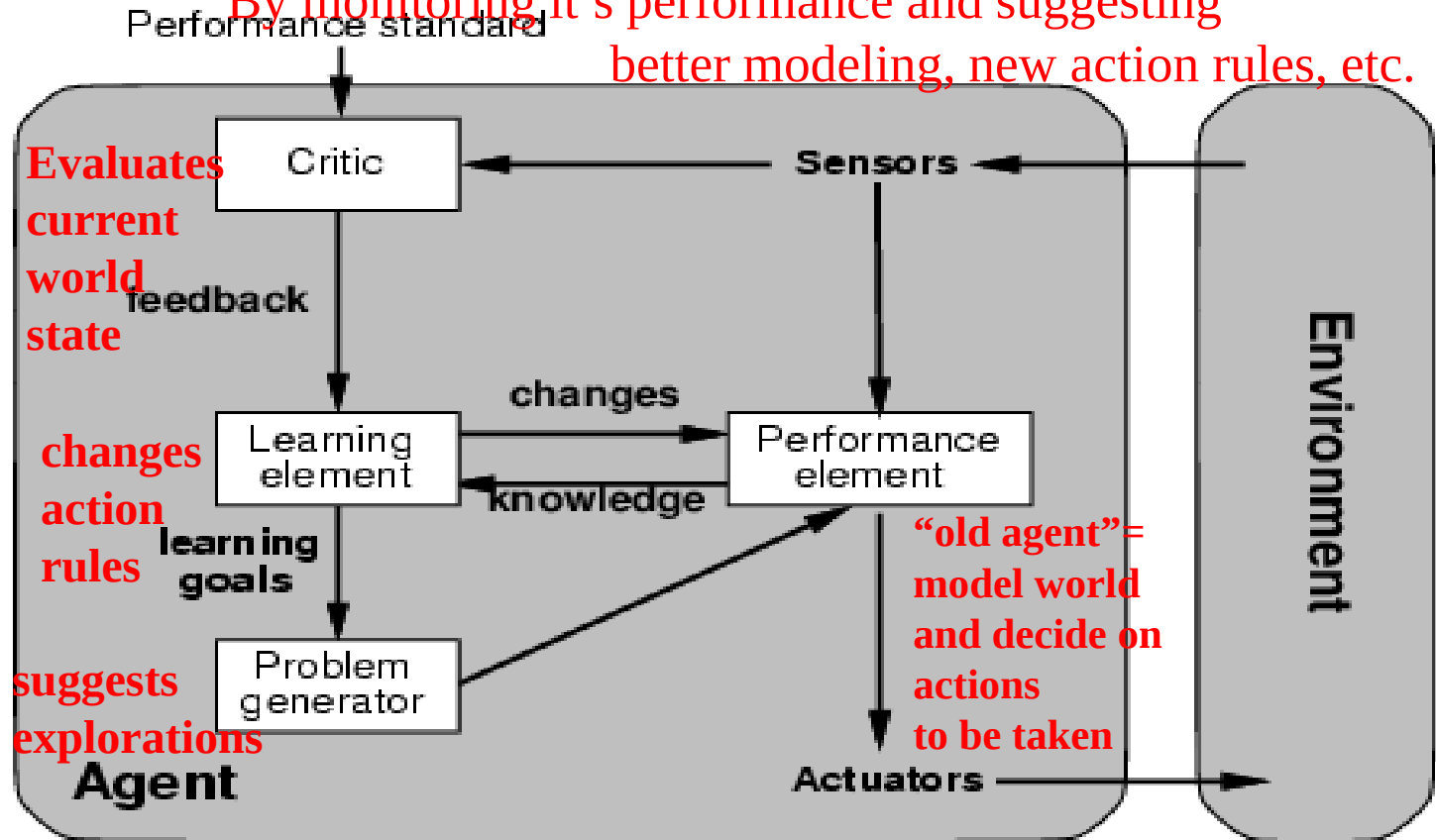
Utility-based agents



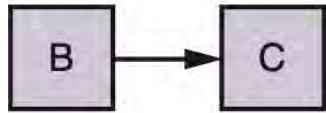
Learning agents

How does an agent improve over time?

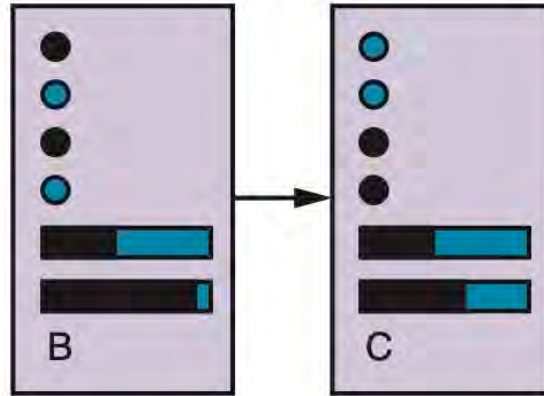
By monitoring its performance and suggesting better modeling, new action rules, etc.



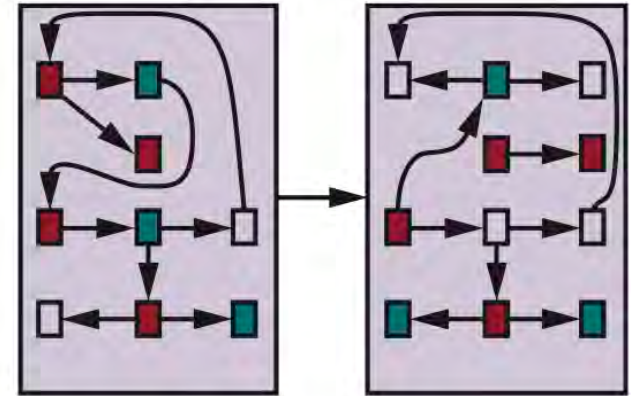
Spectrum of representations



(a) Atomic



(b) Factored



(c) Structured

Summary

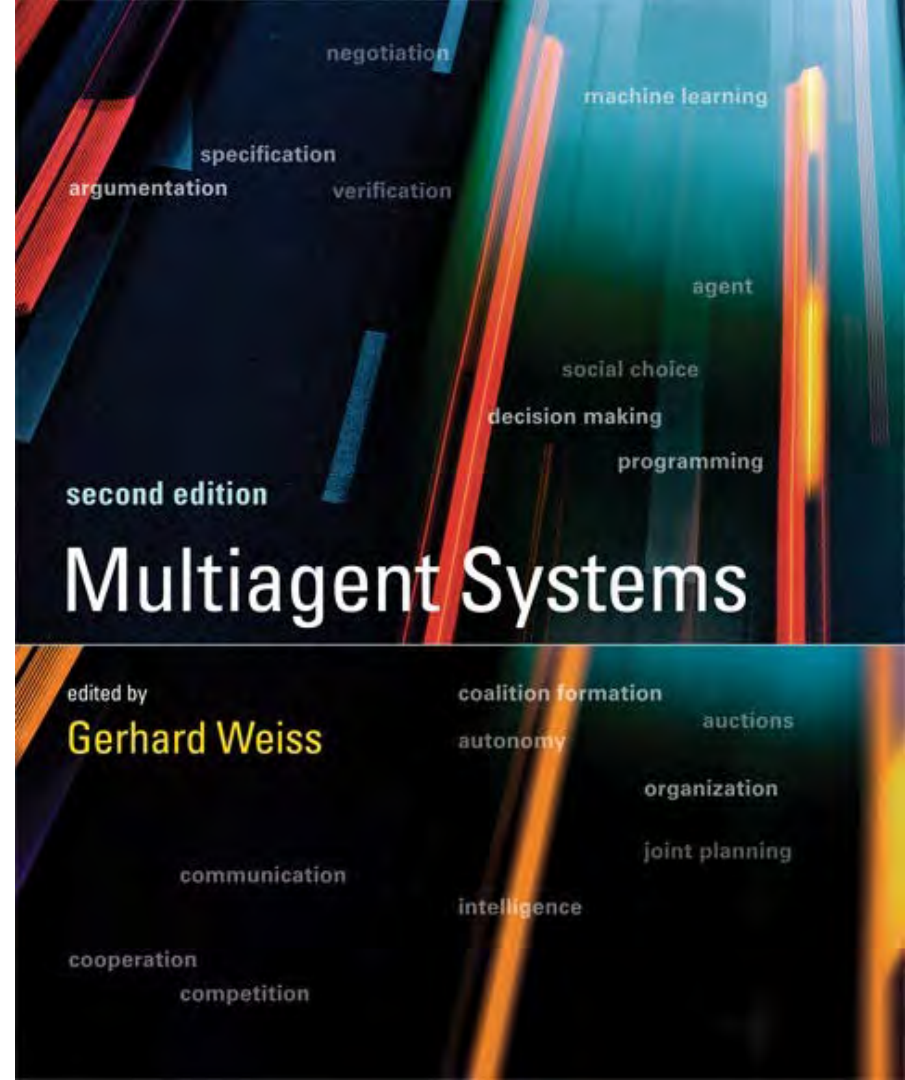
- An *agent* interacts with an *environment* through *sensors* and *actuators*
- The *agent function*, implemented by an *agent program* running on a *machine*, describes what the agent does in all circumstances
- Rational agents choose actions that maximize their expected utility
- PEAS descriptions define task environments; precise PEAS specifications are essential and strongly influence agent designs
- More difficult environments require more complex agent designs and more sophisticated representations

Multiagent Systems

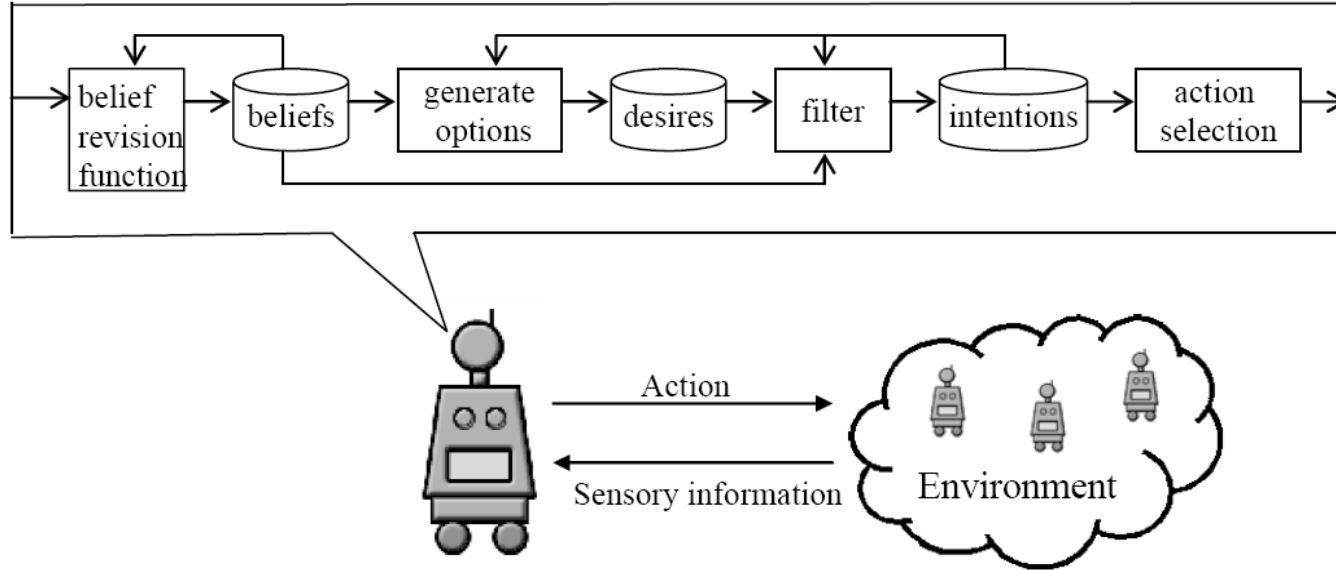
MIT Press

October 2016

[Gerhard Weiss](#)

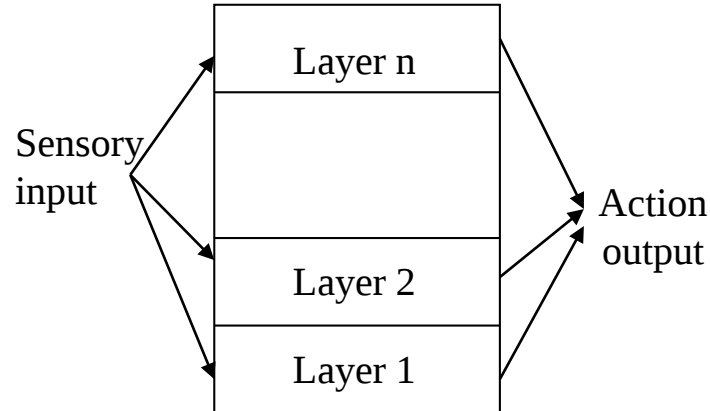


BDI components

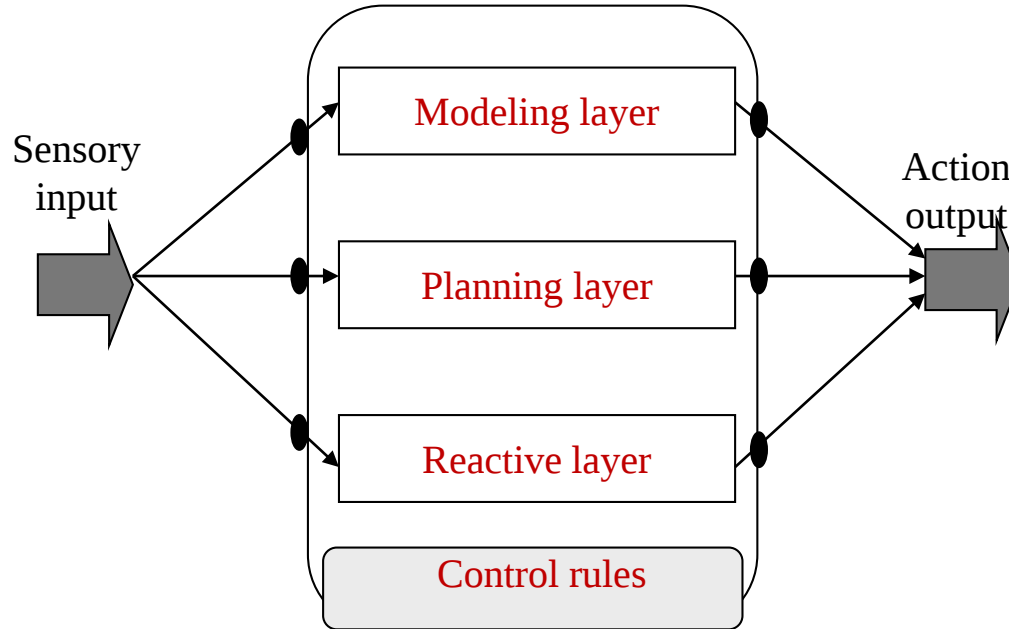


Horizontal layering

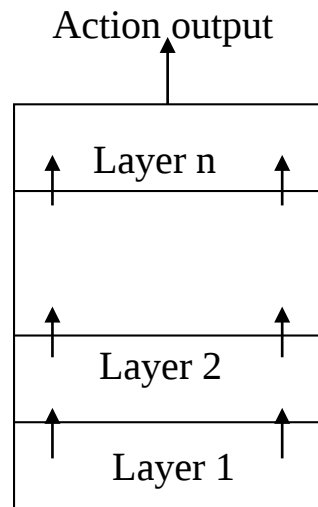
- Each layer can act as an independent agent
- For n different behaviours, n layers are implemented
- The layers compete with each other *in order to take control of the agent*; a mediator function can be introduced



Touring Machines

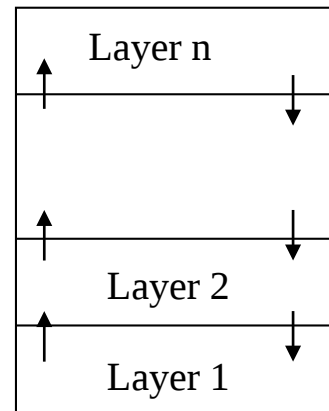


Vertical layering



Sensory input

One-pass

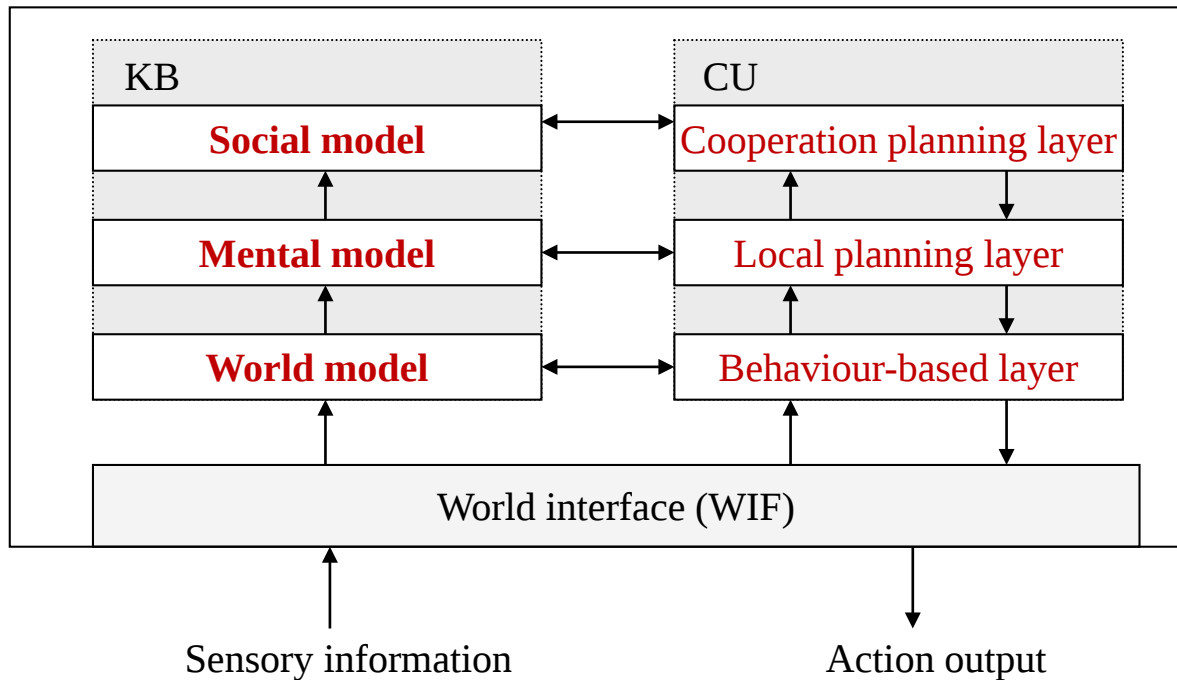


Sensory
input

Action
output

Two-pass

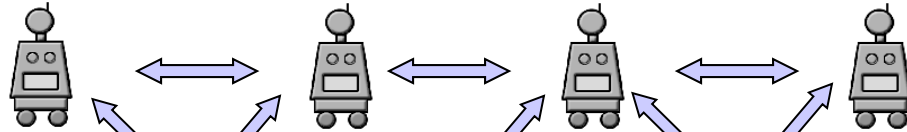
InteRRaP



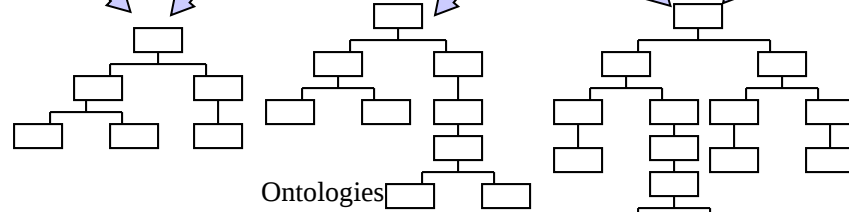
Users



Agents

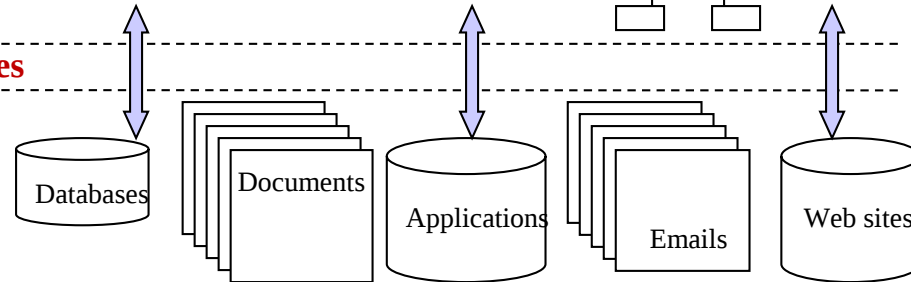


Semantic Web



Web Services

Web



negotiation

machine learning

specification

argumentation

verification

agent

social choice

decision making

programming

second edition

Multiagent Systems

edited by

Gerhard Weiss

coalition formation

auctions

autonomy

organization

joint planning

communication

intelligence

cooperation

competition

Agent Technology for e-Commerce



Maria Fasli



با سپاس از توجه شما دانشجوی گرامی