

Today:

- Examinable Material
- Alternative approaches to AI

Recall:

- Recursive algorithms
- State space, state graph with actions relating states
- Problem = initial state, successor function, cost, goal test
- Tree search strategies (complete, complexity time/space, optimal?)
- Uninformed search: breadth first, depth first, iterative deepening
- Informed/heuristic search: best-first/greedy
- A* search, admissible heuristic, optimality of A*

Recall:

- AI as an attempt to construct artificial intelligent agents.
- Definitions of intelligence
- Systems that **think** or **act rationally** or **like humans**
- Artificial agents: autonomy, interaction, goal direction . . .
- Intentional stance (Dennett)
- Performance, Environment, Actuators, Sensors
- Worst case time complexity, “big O” notation

Recall:

- Local search, hill climbing, simulated annealing
- Constraint solving problems, backtracking search, some heuristics
- Constraint propagation for binary constraints
- Propositional satisfiability as CSP
- Adversarial search and game playing (deterministic, perfect information?)
- Minimax, $\alpha-\beta$ search algorithms, chance moves
- Logical agents, inference about the world
- Propositional logic (syntax, semantics, inference rule, algorithm)

Recall:

- Horn/definite clauses, forward and backward chaining
- Generalised MP, forward & backward chaining in propositional logic
- Distributed execution (definition of linear speed-up, . . .)

Recall:

Planning:

Give a logical description of the **actions**, the **pre-conditions** necessary for the actions to be possible and the **post-conditions** that hold after the action is over.

From a description of the initial state, figure out how to combine actions to achieve a **goal**:
POP algorithm.

Look at R&N, chapter 11.

Recall:

Philosophical Issues:

- Philosophical issues: dualism vs materialism, free will & autonomous agents, mental and physical states
- weak vs strong AI claims
- Searle's Chinese room thought experiment and artificial understanding
- Non-computability in the brain? (Penrose)

Dealing with uncertainty

A lot of knowledge is uncertain, so FOL is not the best mechanism there. We saw incorporation of chance elements in analysing game-playing with non-determinism.

Use ideas from probability – eg Bayesian networks.

Need to know how the probability of an event is affected by finding out about a related event.

eg medical diagnosis — what is the probability of having a disease, given some observed symptoms?

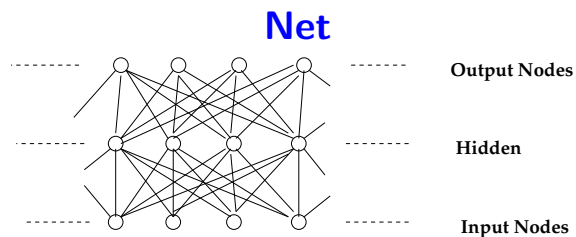
Learning

Many interesting processes can be analysed as a linear sequence of discrete events, influenced by the most recent events.

We can model this by working out probabilities that a particular event will happen, given the last n observed events, and we can **learn** the probabilities by looking at the evolution of the system.

There is a whole area of **Machine Learning**, some of it using symbolic means.

From a general KB, we can deduce facts; try the other direction: from an incomplete KB and some observed facts, work out missing rules in the KB. (inductive Logic Programming)



Net is **trained** by feedback.

- Given input, evaluate the output.
- Reward if it contributes to good output.
- Penalise if contributes to bad output.

After working on a training set, the net can then work on unseen examples.

Alternative Approaches: Connectionism

aka Neural Nets

aka Parallel Distributed Processing.

- Build a system that learns by interaction with the world.
- (loosely) Inspired by the working of the brain.
- Uses a lot of small simple processors.
- Parallel computation (in principle).
- Network signals pass between nodes with varying strengths.

Applications

- Face recognition
- Financial predication
- Sound analysis

The approach seems to work well on tasks where explicit rules are hard to obtain.

We find **graceful degradation** – if a part of the network is removed, the network can still function (though not so well).

KR and Connectionism

Suppose a connectionist system outputs a suggested transaction:

"Sell BA shares"

What happens if we want an explanation?

Usually no explanation is forthcoming
– it's like consulting an oracle.

We can perhaps be convinced by success that we should trust such a system . . .

Hybrid Systems

There is much interest in systems that combine symbolic and sub-symbolic computation.

Part of the net can be

— seen
— programmed)

in a declarative/rule-based way.

KR Hypothesis

Recall: for intelligence, we **must** have a linguistic component that

- represents
- guides behaviour

Spoken Language

Sometimes, patterns in the net suggest an analysis.

Example

recognition of spoken language:

in a trained net, it can be seen that one group of nodes is activated when a vowel is processed, and another when it is a consonant.

So the net has "discovered" that this distinction is important.

A rule can be based on this:

```
if vowel then ....
else ....
```

Parallel Distributed Systems

How do neural nets relate to the KR Hypothesis?

Distributed Representation

- Possible in principle using the KR Hypothesis
- Difficult to see the net as a symbol, or combination of symbols.

Parallel Inference

- Again, possible using the KR Hypothesis.
- In general, no corresponding inference system.

Note that back propagation performs **supposition** rather than deductive **inference**.

A possibility

So it seems that non-hybrid neural nets are ruled out by the KR Hypothesis as intelligent.

A possibility

Scientists like to believe that there are **simple** and **graspable** underlying principles that describe the world.

Logic-based AI wants to explain intelligent action using logic to state such principles.

Combining Approaches

- Often, there is a cultural gap between 2 approaches to AI.
- But, there are many aspects to intelligent behaviour; compare
 - playing chess
 - listening to music

Both approaches can help in the AI Enterprise.

No general rules?

What if there are no such principles?

Maybe intelligence emerges from sufficiently complex unintelligent, unstructured behaviour, **without** any overall organisation principles.

If so, neural nets may be the only way to build a successful AI.

Summary

The End!