

Factored Representations

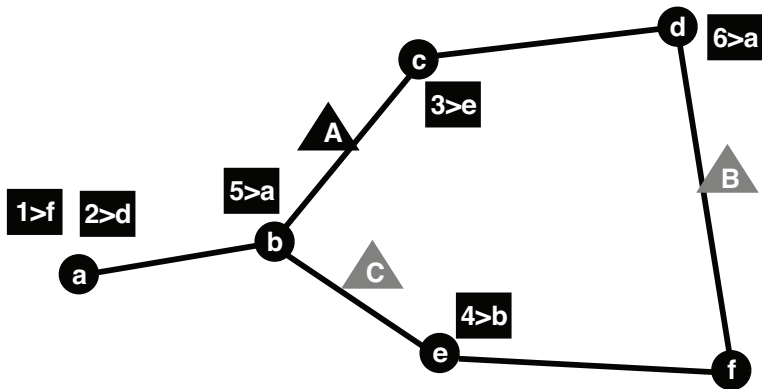
Boi Faltings

Laboratoire d'Intelligence Artificielle

`boi.faltings@epfl.ch`

`http://moodle.epfl.ch/`

Recall: delivery problem



Agent A delivers packages 1..6 to their destinations

Similar real life planning problems

Airport ground-traffic control

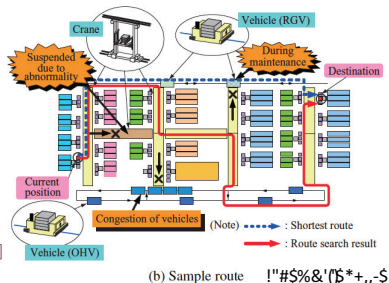
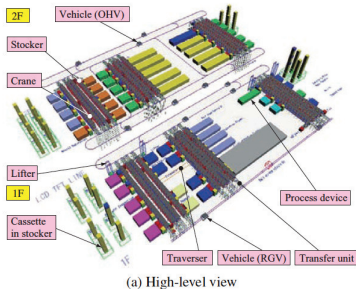
- Guide aircrafts between runway and terminals
- Keep safe distances, minimize taxi and wait times



More real life planning problems

Material control system for LCD manufacturing

- Transfer LCD cassettes to different locations in the plant
- Handle hardware failure and minimize delays



State space explosion

Multiple features $\Rightarrow$ combinatorial explosion of state space:

$\text{pos}(A)=a/b/c/d/e/f$	$\text{pos}(1) = a/b/c/d/e/f$	$\text{holding}(1)=t/f$
	$\text{pos}(2) = a/b/c/d/e/f$	$\text{holding}(3)=t/f$
	...	...
	$\text{pos}(6) = a/b/c/d/e/f$	$\text{holding}(6)=t/f$

$\Rightarrow$ 17'915'904 states

Only few states are reachable

Successor states of

$\{pos(A) = a, pos(1) = a, pos(2) = a, pos(3) = c, \dots\}$:

- $\{pos(A) = a, \dots, holding(1), \dots\}$
- $\{pos(A) = a, \dots, holding(2), \dots\}$
- $\{pos(A) = b, \dots\}$

few actions $\Rightarrow$ few successors.

$\Rightarrow$ construct states dynamically as combination of features.

Many states are equivalent

Can agent move package 1 from a to b:

- $\text{pos}(A) = a/\text{elsewhere?}$
- $\text{pos}(1) = a/\text{elsewhere?}$
- $\text{holding}(1)?$

Differences in other features $\Rightarrow$ equivalent states.

Only $2^3 = 8$ possibilities to distinguish!

$\Rightarrow$ drop features that don't matter!

Factored representations

Idea:

- model each feature by separate predicates (factors).
- represent only those that are important for the current goal.

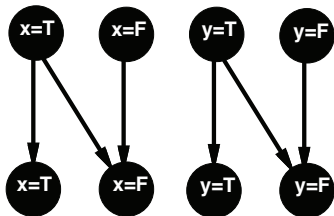
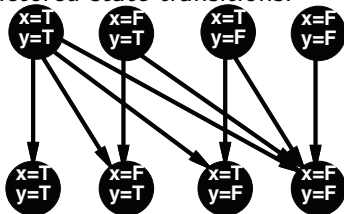
Construct states as needed for planning.

⇒ equivalent states treated as one in the planning process.

⇒ dramatic reduction in complexity.

Factored representations (2)

Factored state transitions:



Treat factors independently $\Rightarrow$ reduce combinatorics.

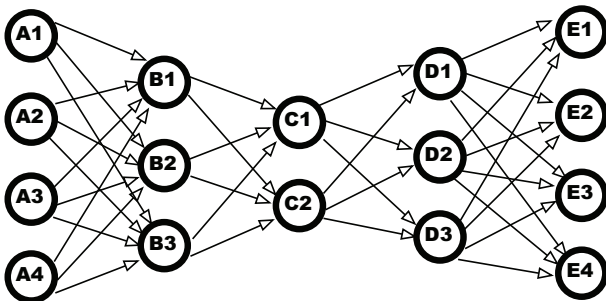
Techniques for factored representations

- State X = vector of k state variables $(x_1, x_2, \dots, x_k)$.
- State variables are also called *fluents*.
- Every combination of state variable values is a state.
- Formulate successor functions and rewards as functions on the vector of state variables.

Logic-based factored representations

- Basis: predicate logic.
- Each state is modelled by a set of true/false predicates.
- De-facto standard: *situation calculus*.
- Most suitable for deliberative agents (planning).

Neural Networks



- Neurons in each layer represent state variables.
- Train autoencoder net to discover a compact factorization:
 $A/E = \text{"raw" state variables}$, $C = \text{factorization (embedding)}$.

Factored Decision Processes

- State $X = (x_1, \dots, x_k)$
- Each x_i is chosen from a finite domain d_i .
- All combinations are allowed $\Rightarrow$ exponential state space.
- Delivery problem: positions of agents, packages, goals.
- How do we express transitions, rewards and policies in a factored way?

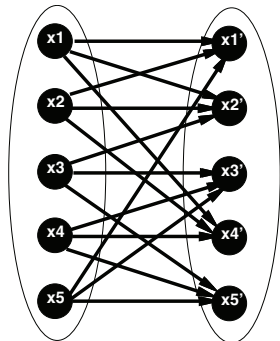
Bayesian Networks

Bayesian network =

- nodes = events, for example $x_3 = c$.
- arcs = causation, for example $(x_3 = c) \rightarrow (x_7 = d)$.
- causation is uncertain, expressed by probability distribution:
for arc $(x_i \rightarrow x_j)$, express $P(x_j|x_i)$ for all $x_j \in d_j, x_i \in d_i$.

Dynamic Bayesian Networks

- Nodes = $(x_1, \dots, x_k) \cup (x'_1, \dots, x'_k)$
- X = states at time t
- X' = states at time $t+1$
- No arcs between x_i s or x'_i s
- Arc from x_i to x'_j if x_i influences x'_j



Model transitions by a separate DBN for each action
(more efficient representations possible)

Probability distributions

Every variable x'_i with inbound nodes $x_j, \dots, x_k$ has a probability distribution:

$$p_i(x'_i | x_j, \dots, x_k)$$

Transition probability between states given as product of distributions:

$$\begin{aligned} & pr(x'_1 = v_1, \dots, x'_n = v_n | x_1 = w_1, \dots, x_n = w_n) \\ &= \prod_{l=1}^n p_l(x'_l = v_l | x_j = w_j, \dots, x_k = w_k) \end{aligned}$$

Best if relations are as local as possible!

Example: Delivery Problem

DBN for agent A moving package 1 from a to b

- $x_1 = \text{pos}(1)$, $x_2 = \text{pos}(A)$, ...
- x'_1 influenced by x_1 and $x_2 \Rightarrow$ arcs
- Probability distributions:

$$p(x'_1 = y | x_1 = x, x_2 = x) = 1, \quad x = a, y = b$$

$$p(x'_1 = y | x_1 = x, x_2 = x) = 0, \quad x \neq a, y \neq b$$

same for $p(x'_2 = y | x_1 = x, x_2 = x)$

Frame axioms:

$$p(x'_i = x | x_i = x) = 1 \quad i \neq 1, 2$$

$$p(x'_i = x | x_i = y) = 0 \quad x \neq y, i \neq 1, 2$$

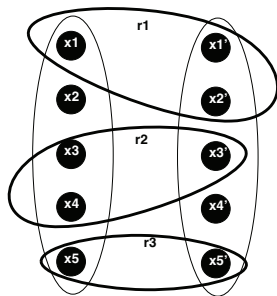
Factoring the policy

- Policy should be formulated on factored state representation.
 - Generally, number of possible actions is small $\Rightarrow$ many states have the same optimal action.
- $\Rightarrow$ Policy = mapping (optimal action) $\Rightarrow$ set of states.
- Assumption: states with same optimal action have a compact description.
 - Examples: combinations of features, decision tree on features.
 - Application of policy: test which action satisfies the current state.

Factoring the reward function

- Instantaneous rewards depend on action and state variables.
- Factored reward function: reward depends on action and *subset* of state variables.
- e.g. delivery problem: agent A moves package 1 from a to b: reward = 1 if successful (as in transition) and b is the goal.
- $\Rightarrow$ dependency on: $\text{pos}(1)$, $\text{pos}(A)$, $\text{destination}(1)$
- Separate reward function for each possible action.
- Select depending on action actually executed.

- Reward functions depend on nodes in dynamic Bayes net.
- Total reward = sum of rewards.
- Can we also factor the value of a state?



Factoring the value function

- Factored reward function $\nRightarrow$ factored value function!
- Value of $\text{pos}(A)=b$ also depends on positions of agents and other packages; if also $\text{pos}(B)=b$ then value increment is lower because agent B can also carry the package from b.
- Factor value function into basis functions $b_i(U)$, $U \subset X$, such that $V(X) = \sum w_i b_i$
- Optimal policy is computed using the same factoring as the value function.

Value function factored into basis functions

- $V(X) = \sum w_i b_i$
- Basis functions b_i are programmed by analysis of the system (ex deliveries: analyze dependencies).
- Weights w_i are determined so that the overall mean square error is minimized.
- Disadvantage: approximation depends on the quality of the basis functions.
- Advantage: often just few basis functions $\Rightarrow$ manageable complexity.

Value iteration?

- As value function $V(S) = \sum w_i b_i(S)$, value iteration recurrence becomes:

$$\begin{aligned} V(S) &= R(S, \pi(S)) + \gamma \sum_{S' \in \mathcal{S}} T(S, \pi(S), S') V(S') \\ \sum w_i b_i(S) &= R(S, \pi(S)) + \gamma \sum_{S' \in \mathcal{S}} T(S, \pi(S), S') \sum w_i b_i(S') \end{aligned}$$

with one equation for each state.

- Many more states than basis functions $\Rightarrow$ more equations than unknowns w_i .
- $\Rightarrow$ solve approximately for least-squares approximation of b_i .
- However, far too many states: intractable to solve.

Policy iteration

Factored policy $\Rightarrow$ factored state space; use for approximating value function in policy iteration:

choose an arbitrary policy π'

loop

$\pi \leftarrow \pi'$

compute the value function V_π of policy π :

solve the linear equations:

$$V_\pi(S) = R(S, \pi(S)) + \gamma \sum_{S' \in \mathcal{S}} T(S, \pi(S), S') V_\pi(S')$$

improve the policy at each state:

$$\pi'(s) \leftarrow \underset{a}{\operatorname{argmax}} (R(s, a) + \gamma \sum_{s' \in \mathcal{S}} T(s, a, s') V_\pi(s'))$$

until $\pi = \pi'$

Alternative: Deep Reinforcement Learning

Neural networks can be seen as a kind of factored representation:

- input units = factors of state representation
(state = *combination* of activation levels)
- output units = factors of action representation.

Deep learning finds a complex policy mapping $state \Rightarrow action$.

Factors of reward function a good starting point for neural units.

Factored representations for deliberative agents

- Factoring helps to reduce complexity of decision processes.
- But still generates policies for all imaginable states $\Rightarrow$ unnecessary complexity.
- Often, need to use deliberative agents.

A simple delivery world

Consider a robot moving packages among a network of locations.
We use the following predicates to model the world in state S :

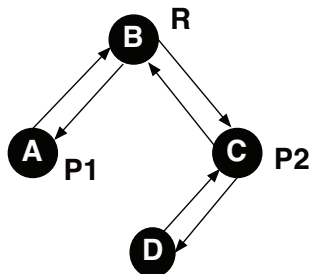
- $AT(p, l, S)$: object p is at location l
- $LOC(l, S)$: the robot is at location l
- $CONN(l1, l2, S)$: there is a connection from location $l1$ to $l2$

Argument S allows different truth values in different states.

Modeling world states

States modelled by a set of propositions, e.g. state S_0 :

$AT(P1, A, S_0)$
 $AT(P2, C, S_0)$
 $LOC(B, S_0)$
 $CONN(A, B, S_0)$
 $CONN(B, A, S_0)$
 $CONN(B, C, S_0)$
 $CONN(C, B, S_0)$
 $CONN(C, D, S_0)$
 $CONN(D, C, S_0)$



Situations and states

States contain unnecessary details:

$AT(P2, C, S_0)$ *not important for plan* $\Rightarrow$ *drop*

$\Rightarrow$ partial models = *Situations*

- Situations allow specifying goals, more general plans.
- Operators model agent actions.
- In *situation calculus*, operators transform situations.

Operators (situation calculus = STRIPS)

An operator $S_i \Rightarrow S_{i+1}$ is characterized by:

- *preconditions*: propositions which must be true in S_i for the operator to be applicable.
- *postconditions*: propositions which will be true in S_{i+1} as a result of the action. Also called ADD-LIST.
- *DELETE-LIST*: propositions which will no longer be true in S_{i+1} . Often identical with preconditions.

Formalized in the *Planning Domain Definition Language* (PDDL).

Example: CARRY

$\text{CARRY}(p, l1, l2, S_i)$ models the action of the robot carrying object p from location $l1$ to $l2$.

It is defined as follows:

- preconditions (P): $\text{CONN}(l1, l2, S_i)$, $\text{LOC}(l1, S_i)$, $\text{AT}(p, l1, S_i)$
- add-list (A): $\text{LOC}(l2, S_{i+1})$, $\text{AT}(p, l2, S_{i+1})$
- delete-list (D): $\text{LOC}(l1, S_{i+1})$, $\text{AT}(p, l1, S_{i+1})$

Example: MOVE

$\text{MOVE}(11, 12, S_i)$ models the action of the robot moving from location 11 to 12 without carrying anything.

It is defined as follows:

- preconditions (P): $\text{CONN}(11, 12, S_i)$, $\text{LOC}(11, S_i)$
- add-list (A): $\text{LOC}(12, S_{i+1})$
- delete-list (D): $\text{LOC}(11, S_{i+1})$

Situation calculus as a factored representation

- Most factors of the world state have no influence on current action.
- For example, moving a package in a delivery problem is not influenced by the positions of other packages.
- Situation calculus does not represent these
⇒ automatically groups equivalent states.
- Situation calculus operators model progression of situations as a result of actions, similar to dynamic Bayes nets.

Least commitment

- Plans are sequences of multiple actions: can we factor as well?
 - Assume actions $A, \dots, Z$ can be carried out in any order.
- ⇒ 26! different (ordered) plans:

$$(A, B, \dots, Z), (B, A, \dots, Z), \dots$$

Search treats all of them as separate alternatives!

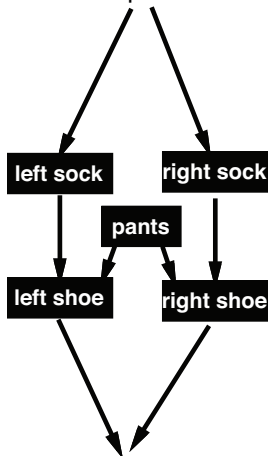
- Delay commitment on order ⇒
1 plan with 26 parallel actions:

$$(\dots, \{A, B, \dots, Z\}, \dots)$$

Less to enumerate ⇒ more efficient search.

Non-linear (partial-order) planning

Plan indicates partial orders:



Boxes = operations

Arrows = precedence relations

Making nonlinear planning efficient

- Problem: complexity of establishing causal links in a non-conflicting fashion.
 - Idea: explicitly generate sets of actions that can be executed in parallel.
- ⇒ plan = sequence of (potentially parallel) action groups.
- Generate the groups to avoid conflicts in causal links.
- ⇒ no need for backtracking on this link structure.

Constructing a finite decision space

- Assumptions:
 - finite universe $\Rightarrow$
situation calculus uses a finite set of propositions
 - plan has at most n steps (use iterative deepening)
- Each proposition becomes a state variable $\Rightarrow$ universe of state variables for each of the n states.
- Only reachable propositions are represented!

Example

5 state variables:

- ls, rs : wearing left/right sock
- lh, rh : wearing left/right shoe
- p : wearing pants

Initial state: $\neg ls, \neg rs, \neg lh, \neg rh, \neg p$

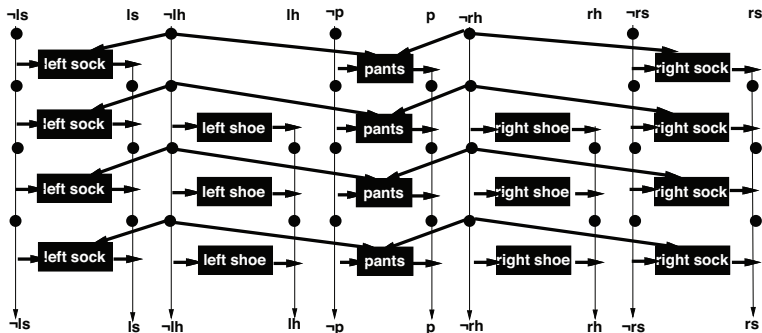
5 operators:

- left/right sock:
 $P = \{\neg ls/rs, \neg lh/rh\}, D = \{\neg ls/rs\}, A = \{ls/rs\}$
- left/right shoe:
 $P = \{\neg lh/rh, ls/rs\}, D = \{\neg lh/rh\}, A = \{lh/rh\}$
- pants:
 $P = \{\neg p, \neg lh, \neg rh\}, D = \{\neg p\}, A = \{p\}$

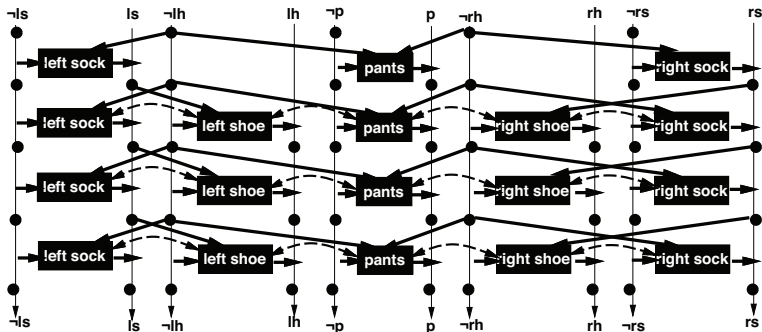
Planning graphs

- Idea: consider what actions can be carried out simultaneously.
- Graph has *layers*. Each layer contains nodes
 - for each possible proposition that could hold
 - for each possible action
- First layer: initial state + actions possible in initial state.
- Following layers: propositions and actions that might be possible.
- Construction: polynomial time.

Planning graph example



Exclusion relations



- Interference, Competing needs = constraints!
- Result: identify what actions can take place in parallel.

Backward search

- Goals at time t: ls, rs, lh, rh, p
- Select set of non-exclusive actions at time t:
 - 1 {pants, left sock, right sock}
 - 2 {left shoe, right shoe}

⇒ Goals at time t-1:

- 1 $\neg p, \neg ls, \neg rs, \neg lh, \neg rh, lh, rh$
 - 2 $\neg lh, \neg rh, ls, rs, p$
- Goals (1) contradictory ⇒ backup.

Search (continued)

- Goals (2) $\Rightarrow$ non-exclusive actions:
 {left sock, right sock, pants}

$\Rightarrow$ Goals at time t-2:
 $\neg p, \neg ls, \neg rs, \neg lh, \neg rh$

Satisfied by initial conditions!

- Plan:
 - 1 {left sock, right sock, pants}
 - 2 {left shoe, right shoe}

Testing for unsolvability

- After n levels, all levels of the graph will become identical.
- If the goal state is not contained in this state, or ruled out by mutual-exclusion constraints, the problem is unsolvable.
- Extension of this criterion for general unsolvability test.

Does Graphplan give all solutions?

No - plan not included:

left sock, pants, left shoe, right sock,
right shoe

However, complete:

no solution in Graphplan

$\Rightarrow$ *no plan exists*

Using Graphplan as a heuristic

- Simplification: ignore all delete lists.
- ⇒ after forming feasible sets, can find a plan without backtracking.
- Use Graphplan on simplified problem as a heuristic for planning with A* algorithm:
Fast Forward algorithm.
 - FF often much faster than Graphplan itself.

Graphplan $\Rightarrow$ Satisfiability

- Consider collection of *clauses*:

$$I1 \vee I2 \vee I3 \vee \dots$$

for example:

$$(AT(P1, A, S_0) \wedge CARRY(P1, A, B, S_0)) \Rightarrow AT(P1, B, S_1)$$

$\Rightarrow$ clause:

$$\neg AT(P1, A, S_0) \vee \neg CARRY(P1, A, B, S_0) \vee AT(P1, B, S_1)$$

- Satisfiability (SAT): what truth assignment will make a collection of clauses simultaneously true.
- Research community has obtained efficient search tools for this problem.

Graphplan as SAT

Variables:

- each operator at each time instant.
- each proposition at each time instant.

Constraints (clauses):

- each operator with its preconditions in preceeding state.
- each operator with its postconditions in following state.
- exclusions among propositions in each state.
- exclusions among operators using the same resource.

Assumptions

Same as for Graphplan:

- finite set of propositions (fluents).
- plan has at most n steps.

Time broken up into $2n$ points:

- even: states.
- uneven: actions.

Encoding as literals

Literals of the SAT problem =

- all propositions describing states at every even time, e.g.:
 $AT(P1, A, 0), AT(P1, A, 2), AT(P1, A, 4), \dots, AT(P1, A, 2n)$
 $AT(P1, B, 0), AT(P2, A, 0), AT(P2, B, 0), \dots$
- all possible actions at every odd time, e.g.:
 $MOVE(A, B, 1), MOVE(A, B, 3), \dots, MOVE(A, B, 2n - 1)$
 $CARRY(P1, A, B, 1), CARRY(P1, A, B, 3), \dots$

Axioms:

- INIT, GOAL: for initial and final conditions.
- for every action op and every odd time t , an implication

$$op(t) \Rightarrow P(t - 1), A(t + 1), \neg D(t + 1)$$

(Notation: $E = A \cup \neg D$)

- frame axioms.

Frame Axioms

- Frame axioms: ensure that propositions not affected by actions remain true.
- Classical frame axioms: associated with every operator

$$AT(P1, C, t - 1) \wedge CARRY(P2, A, B, t) \Rightarrow AT(P1, C, t + 1)$$

Complex: requires a separate axiom for every action and every proposition.

- Explanatory frame axioms:

$$AT(P2, C, t - 1) \wedge \neg AT(P2, C, t + 1) \Rightarrow \\ CARRY(P2, C, D, t) \vee CARRY(P2, C, B, t) \vee \dots$$

- Exclusion constraints: actions with conflict between precondition/effect cannot be executed in parallel.

Comparison with MDP

MDP:

- any state can be successor to another state with some probability.
- $\Rightarrow$ need to immediately plan for all states...

$\Rightarrow$ focus instead on limited uncertainty:

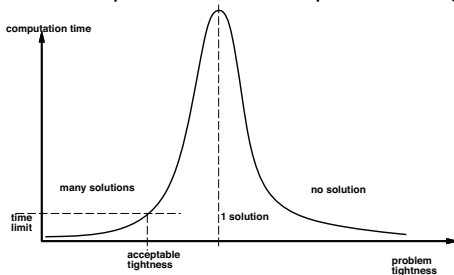
- uncertain effects (but with limited set of choices)
- conditional effects (predictable)
- measurement actions

can be integrated particularly well with SAT and constraint satisfaction techniques.

Active research area.

Phase transition behavior

- SAT expected computation time depends on *tightness*:



⇒ agent has to operate with sufficient liberty.

Summary

- Factored, logical representations can greatly reduce complexity of planning.
- Can improve efficiency of reactive agents (but complex).
- Neural nets can learn factored representations that fit well with reactive agents (deep reinforcement learning).
- Factoring is widely used in deliberative agents.
- Efficiency gains through least-commitment principle:
 - operator order: non-linear, partial-order planning
 - operator choice: graphplan, SATplan