

Automated Planning

Plan-Space Planning / Partial Order Causal Link Planning

Jonas Kvarnström

Automated Planning Group

Department of Computer and Information Science

Linköping University

Partly adapted from slides by Dana Nau

Licence: Creative Commons Attribution-NonCommercial-ShareAlike, <http://creativecommons.org/licenses/by-nc-sa/2.0/>

jonas.kvarnstrom@liu.se – 2018

Motivations

Motivating Problem

- Simple planning problem:

- Two crates

- At A
 - Should be at B



- One robot

- Can carry up to two crates
 - Can move between locations, which requires one unit of fuel
 - Has only two units of fuel



Let's see what a forward-chaining planner *might* do (depending on heuristics)...

Motivating Problem 2: Forward Search

We have:

robotat(A)
at(c1,A)
at(c2,A)
has-fuel(2)

pickup(c1,A)
effects:
holding(c1),
¬at(c1,A)

robotat(A)
holding(c1)
at(c2,A)
has-fuel(2)

drive(A,B)
effects:
¬robotat(A),
robotat(B)

robotat(B)
holding(c1)
at(c2,A)
has-fuel(1)

put(c1, B)
effects:
at(c1,B),
¬holding(c1)

robotat(B)
at(c1,B)
at(c2,A)
has-fuel(1)

pickup(c1, B)
effects:
¬at(c1,B),
holding(c1)

robotat(B)
holding(c1)
at(c2,A)
has-fuel(1)

Cycle, backtrack

drive(B,A)
effects:
¬robotat(B),
robotat(A)

robotat(A)
at(c1,B)
at(c2,A)
has-fuel(0)

pickup(c2,A)
effects:
holding(c2),
¬at(c2,A)

robotat(A)
at(c1,B)
holding(c2)
has-fuel(0)

Dead end,
backtrack

drive(B,A)
effects:
robotat(A),
¬robotat(B)

robotat(A)
holding(c1)
at(c2,A)
has-fuel(0)

Dead end, backtrack

Why is this
not a cycle?

We want:

at(c1,B)
at(c2,B)

Motivating Problem 3

We have:

robotat(A)
at(c1,A)
at(c2,A)
has-fuel(2)

pickup(c1,A)

effects:
holding(c1),
¬at(c1,A)

robotat(A)
holding(c1)
at(c2,A)
has-fuel(2)

drive(A,B)

effects:
¬robotat(A),
robotat(B)

robotat(B)
holding(c1)
at(c2,A)
has-fuel(1)

put(c1, B)

effects:
at(c1,B),
¬holding(c1)

robotat(B)
at(c1,B)
at(c2,A)
has-fuel(1)

Keep
backtracking...

pickup(c2,A)

effects:
holding(c2),
¬at(c2,A)

robotat(A)
holding(c1)
holding(c2)
has-fuel(2)

drive(A,B)

effects:
¬robotat(A),
robotat(B)

robotat(B)
holding(c1)
holding(c2)
has-fuel(1)

put(c1, B)

effects:
at(c1,B),
¬holding(c1)

robotat(B)
at(c1,B)
holding(c2)
has-fuel(1)

put(c2, B)

effects:
at(c2,B),
¬holding(c2)

robotat(B)
at(c1,B)
at(c2,B)
has-fuel(1)

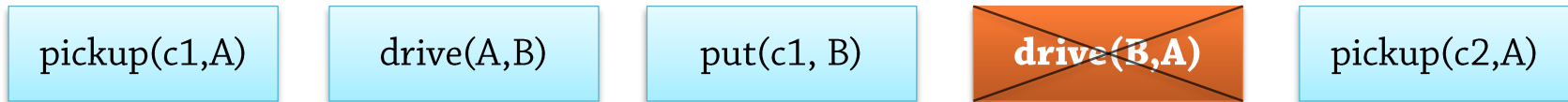
We want:

at(c1,B)
at(c2,B)

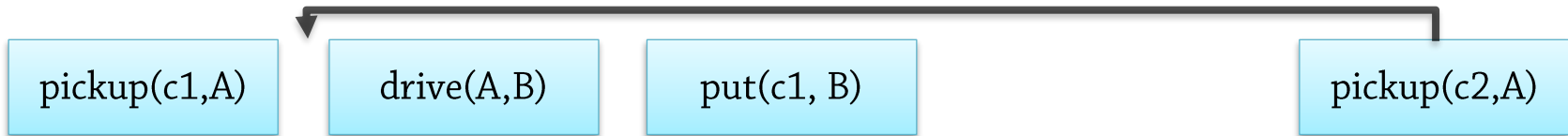
Motivating Problem 4

■ Observations:

- **Most** actions we added before backtracking were **useful** and **necessary**!



- At first, we added them in the **wrong order**



- **Forward** and **backward** planning **commits** immediately to action order
 - Puts each action in its **final place** in the plan
- State space **heuristics** must tell us:
 - Which actions are **useful**
 - ***When to add them to the plan***

What if we could **rearrange** actions?

Planning Outside the State Space

■ Sequences with arbitrary insertion: Useful?

- Add actions in sequence, as in state space planning...

pickup(c1,A)

drive(A,B)

put(c1, B)

- Realize you need another one...

pickup(c2,A)

How to decide which action to *insert*, if not at the end

- Make a space...

pickup(c1,A)

drive(A,B)

put(c1, B)

How to decide where to insert it?

- ...and place the action there

pickup(c1,A)

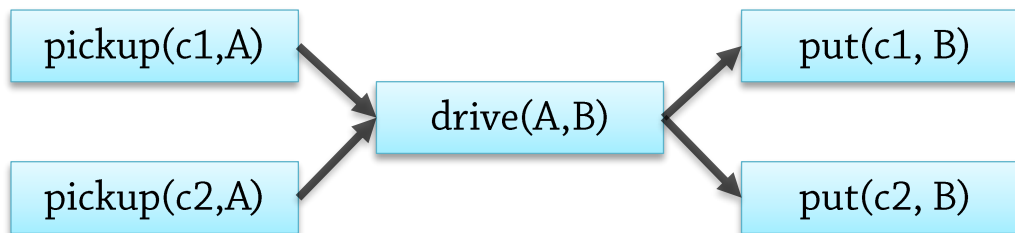
pickup(c2,A)

drive(A,B)

put(c1, B)

How to check that "old" preconditions remain satisfied?

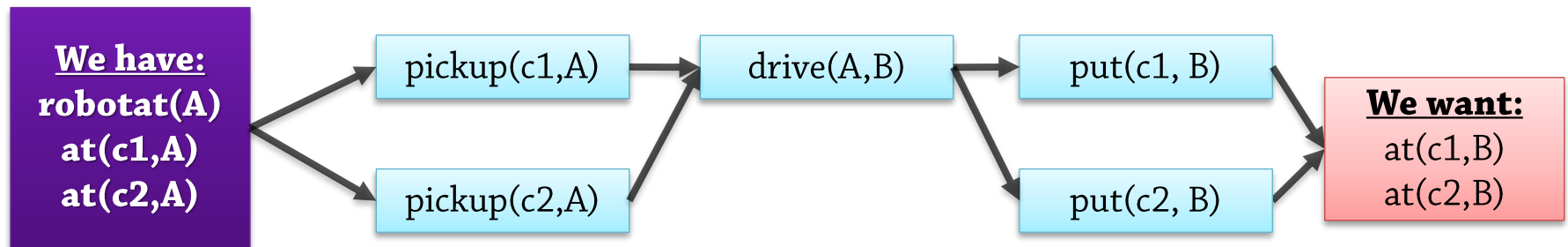
- If we must deal with this complexity:
 - We can "get more for the same price"
- Let's skip sequences completely – a plan could be *partially ordered*:



- A set of **actions** $A = \{ a_1, a_2, a_3, \dots \}$
- A set of **precedence constraints** $\{ a_1 < a_2, a_1 < a_3, \dots \}$
 - a_1 must finish before a_2 starts, ...
 - Here: **solid arrows**

How do we *generate* these plans?

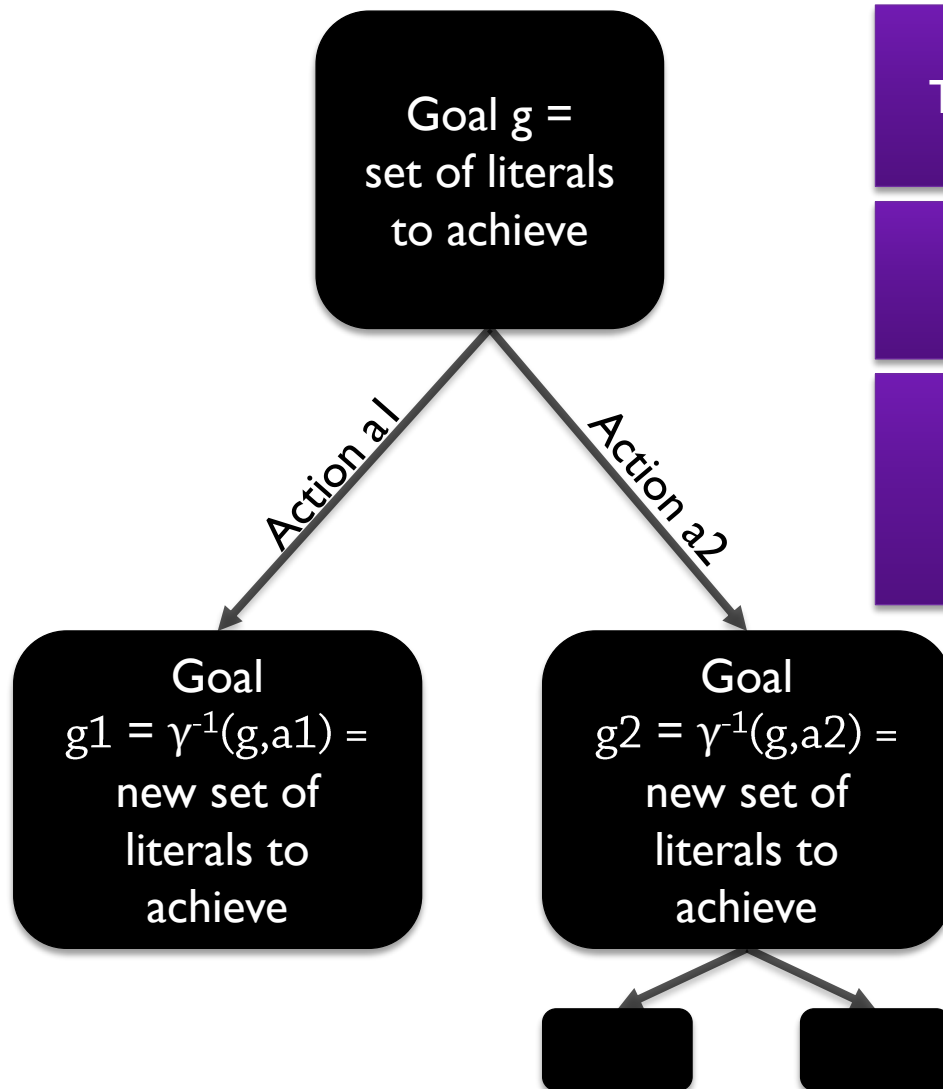
- **Partial Order Causal Link** (POCL) planning
 - Use a partial order, as described
 - Not when executing the plan
 - Only to **delay commitment** to ordering
 - As in backward search:
 - Add **useful** actions to achieve necessary conditions
 - Keep track of what **remains** to be achieved
 - But: Insert actions "**at any point**" in a plan



More sophisticated "bookkeeping" required!

POCL 2: Comparison to Backward Search

- Search tree for backward search, earlier:



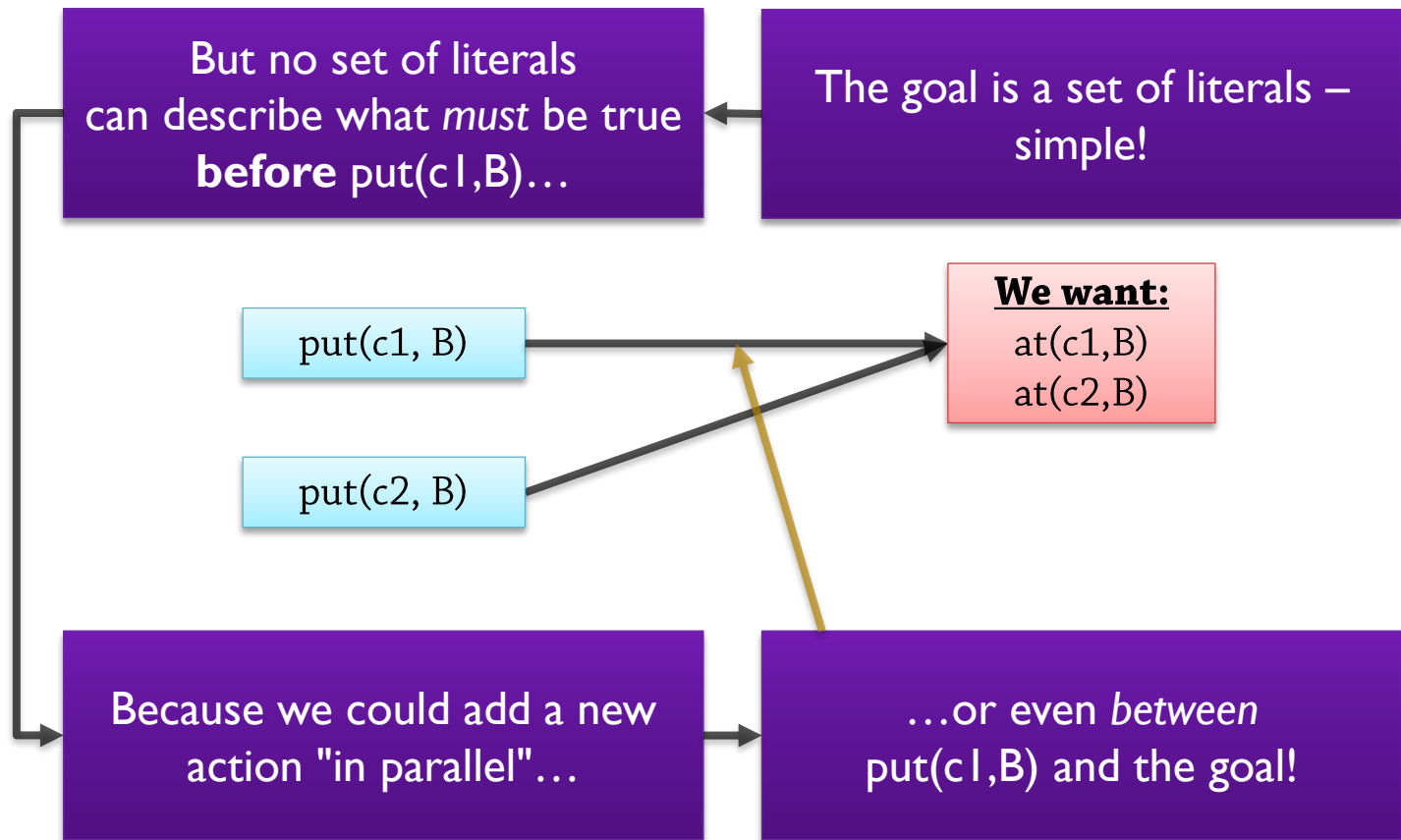
The goal is a set of literals – simple!

Every step takes you to a new set of literals to achieve

From a search node, you know how to reach the goal using a *sequence* of actions

POCL 3: Comparison to Backward Search

- In POCL planning:
 - There is no sequence – and no clear "before" relation!



Has consequences for the POCL plan structure...

POCL 4: Conditions; Goal Action

- Must keep track of individual propositions to be achieved

- Throughout the plan – not a single state $g1 = \gamma^{-1}(g, a1)$

- May come from preconditions of every action in the plan

Notation chosen for this presentation: Preconditions on the left/top side

robotat(B)

holding(c1)

put(c1, B)

- May come from the problem goal as in backward search

- Let's use a uniform representation

- Add a "fake" goal action to every plan, with the goals as preconditions!

at(c1, B)

at(c2, B)

goalaction

preconds: the goal
effects: none

POCL 5: Effects; Initial Action

- Must keep track of individual propositions that are achieved
 - Throughout the plan – not from a single *relevant* action
 - May come from effects of every action in the plan

Notation chosen for this presentation: Effects on the right/bottom side

robotat(B)

holding(c1)

put(c1, B)

at(c1,B)

-holding(c1)

- May come from the initial state
 - Add a "fake" initial action to every plan, with the initial state as effects!
- Effects are sometimes omitted from the slides, due to lack of space...

initaction

preconds: none
effects:
the initial state

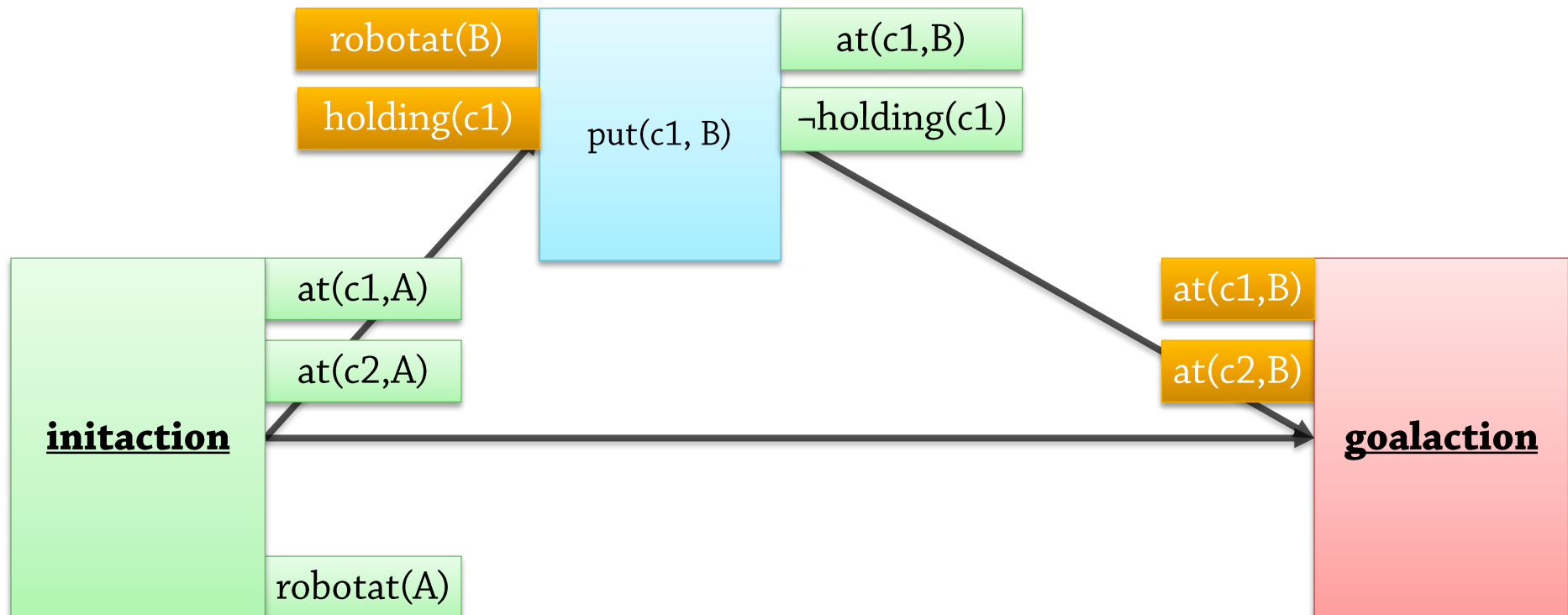
at(c1,A)

at(c2,A)

robotat(A)

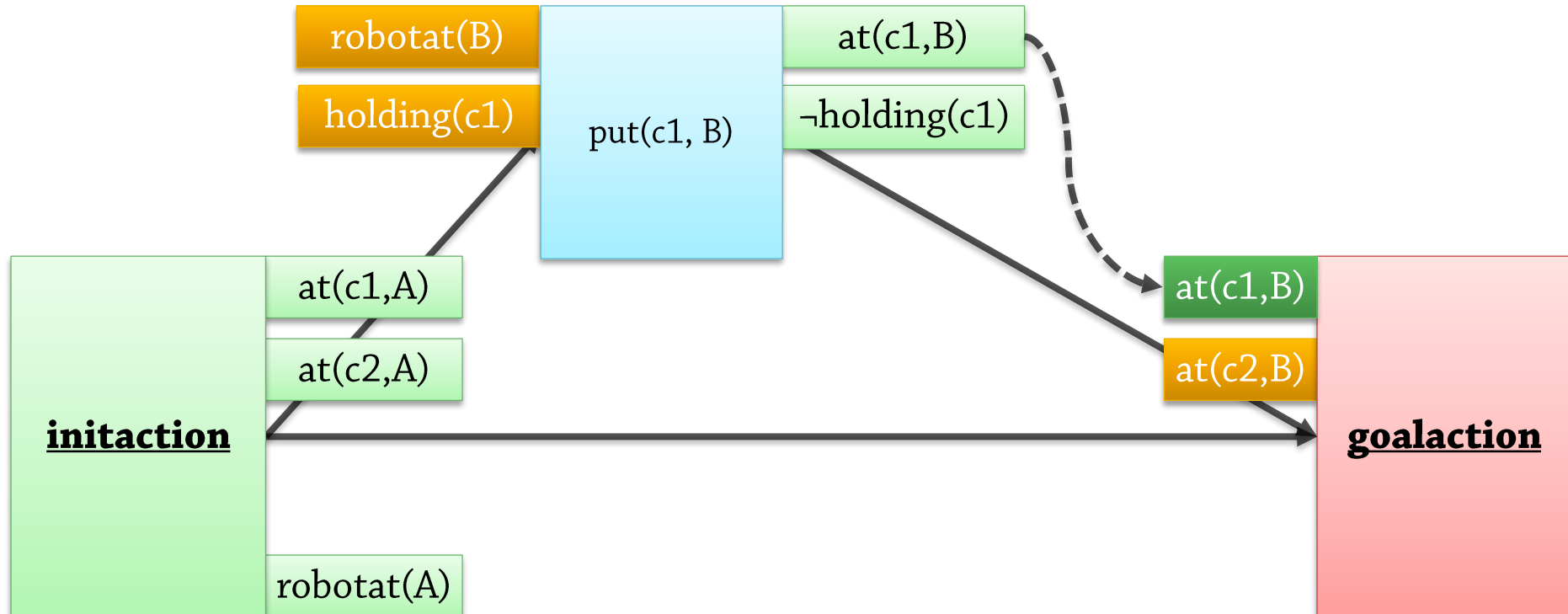
POCL 6: Precedence Constraints

- Plan structure so far:



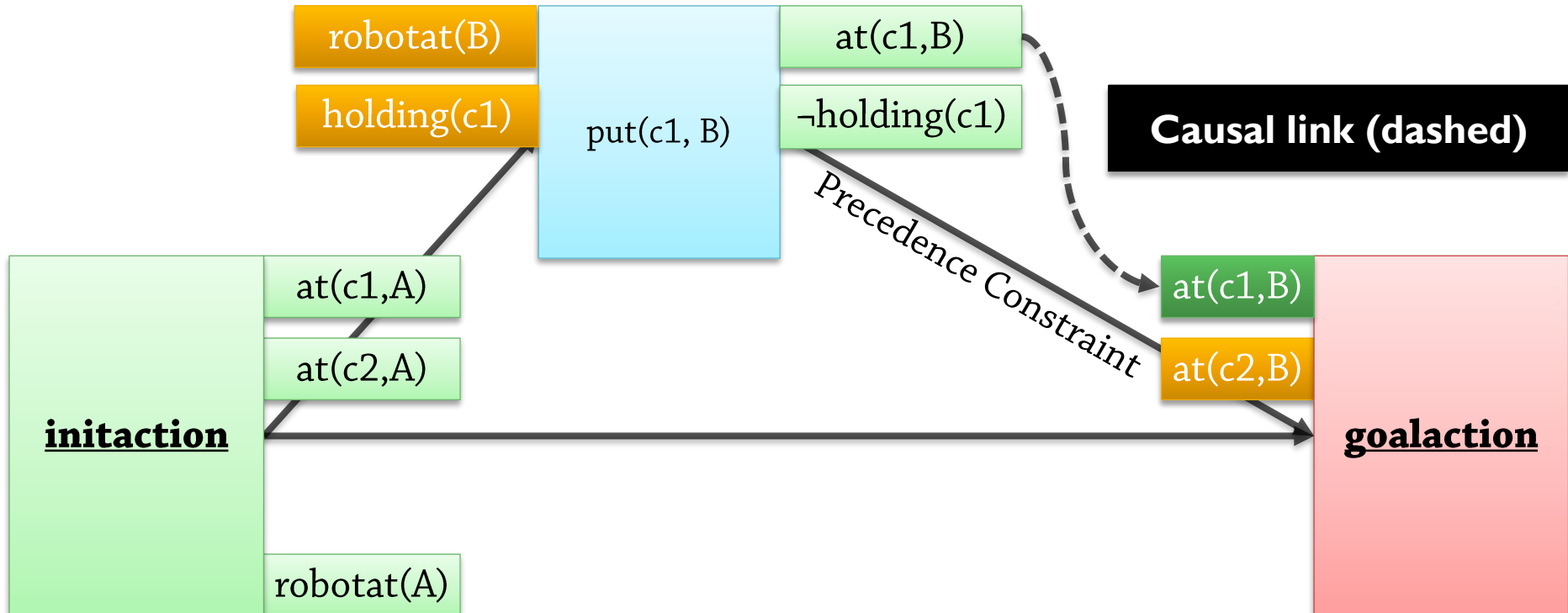
- Must keep track of which action achieves which precondition
 - Causal links

Causal link (dashed):
at(c1,B) must **remain true**
between the end of *put* and the beginning of *goalaction*.
No one must delete it!
Important for *threat management* (later)



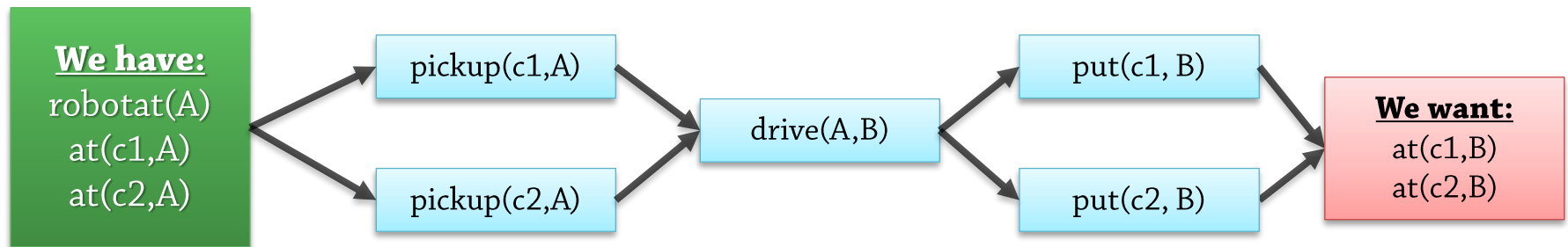
POCL 8: Partial-Order Plans

- To summarize, a ground **partial-order plan** consists of:
 - A set of **actions**
 - A set of **precedence constraints** $a \rightarrow b$
 - Action a must precede b
 - A set of **causal links** $a \xrightarrow{p} b$
 - Action a establishes the precondition p needed by b



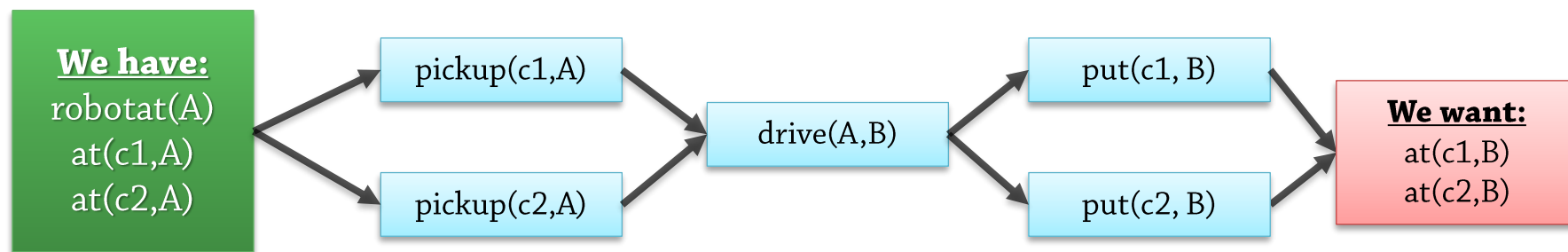
- Original motivation: **performance**

- Therefore, a partial-order plan is a **solution** iff **all sequential** plans satisfying the ordering are solutions
 - Similarly, **executable** iff corresponding sequential plans are executable
 - $\langle \text{pickup}(c1,A), \text{pickup}(c2,A), \text{drive}(A,B), \text{put}(c1,B), \text{put}(c2,B) \rangle$
 - $\langle \text{pickup}(c2,A), \text{pickup}(c1,A), \text{drive}(A,B), \text{put}(c1,B), \text{put}(c2,B) \rangle$
 - $\langle \text{pickup}(c1,A), \text{pickup}(c2,A), \text{drive}(A,B), \text{put}(c2,B), \text{put}(c1,B) \rangle$
 - $\langle \text{pickup}(c2,A), \text{pickup}(c1,A), \text{drive}(A,B), \text{put}(c2,B), \text{put}(c1,B) \rangle$



Partial Orders and Concurrency

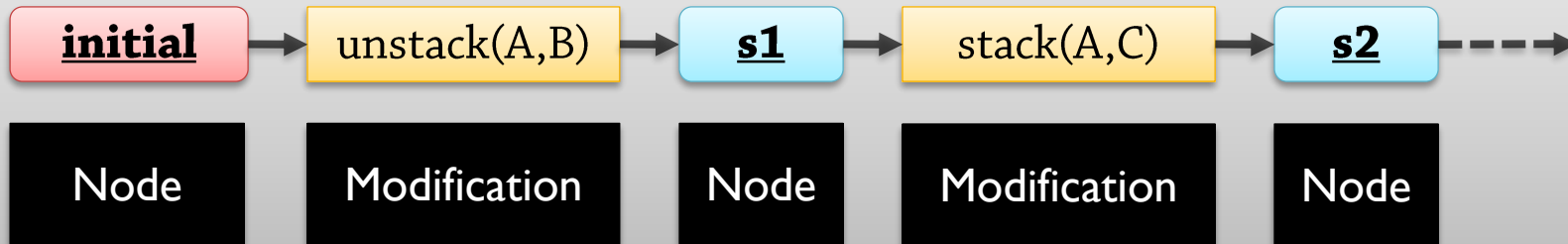
- Can be extended to allow concurrent execution
 - Requires a new formal model!
 - Our state transition model *does not define* what happens if $c1$ and $c2$ are picked up simultaneously!



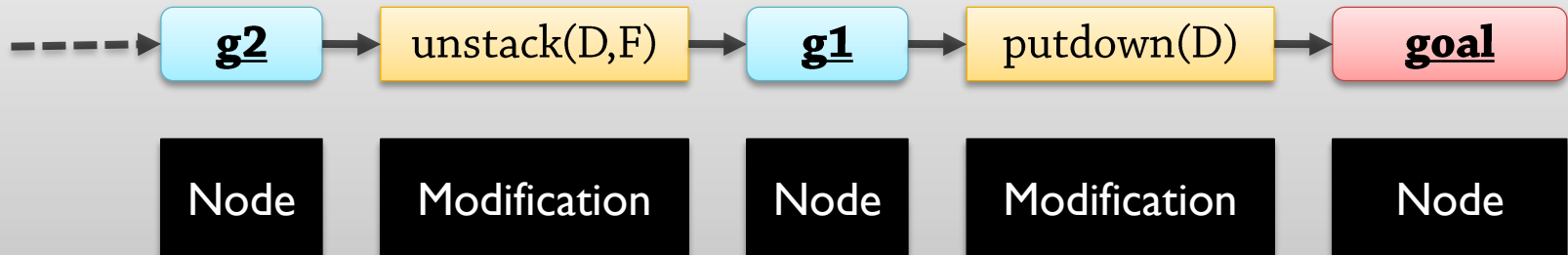
Generating Partial-Order Plans

Context: Forward, Backward

Forward search: A search node is a "current state"

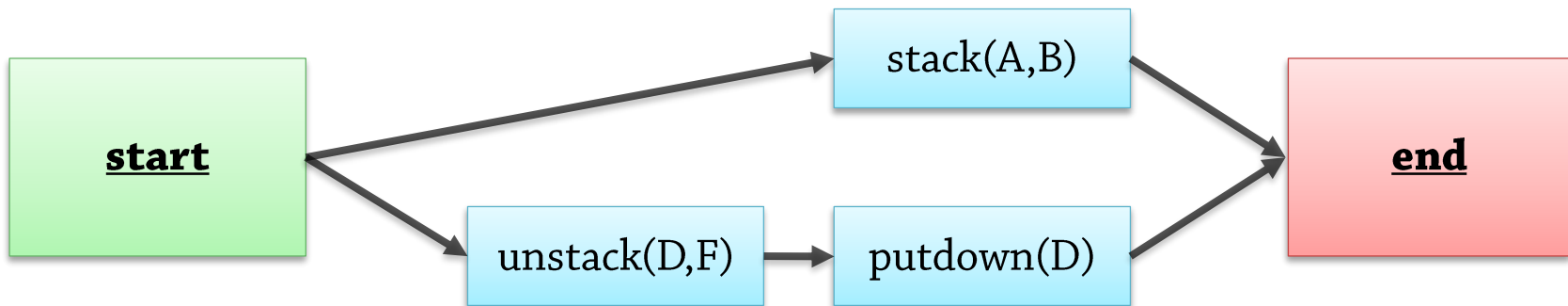


Backward search: A search node is a "current goal"



No Current State during Search!

- With **partial-order plans**: No “current” state or goal!
 - What is true after `stack(A,B)` below?
 - **Depends** on the order in which **other** actions are executed
 - **Changes** if we insert **new** actions **before** `stack(A,B)`!



A search node can't correspond to a state or goal!

Search Nodes are Partial Plans

- A node has to contain more information: The entire plan!
 - The initial search node contains the initial plan
 - The special initial and goal actions
 - A precedence constraint

Therefore, this is
one form of
"plan-space" planning!



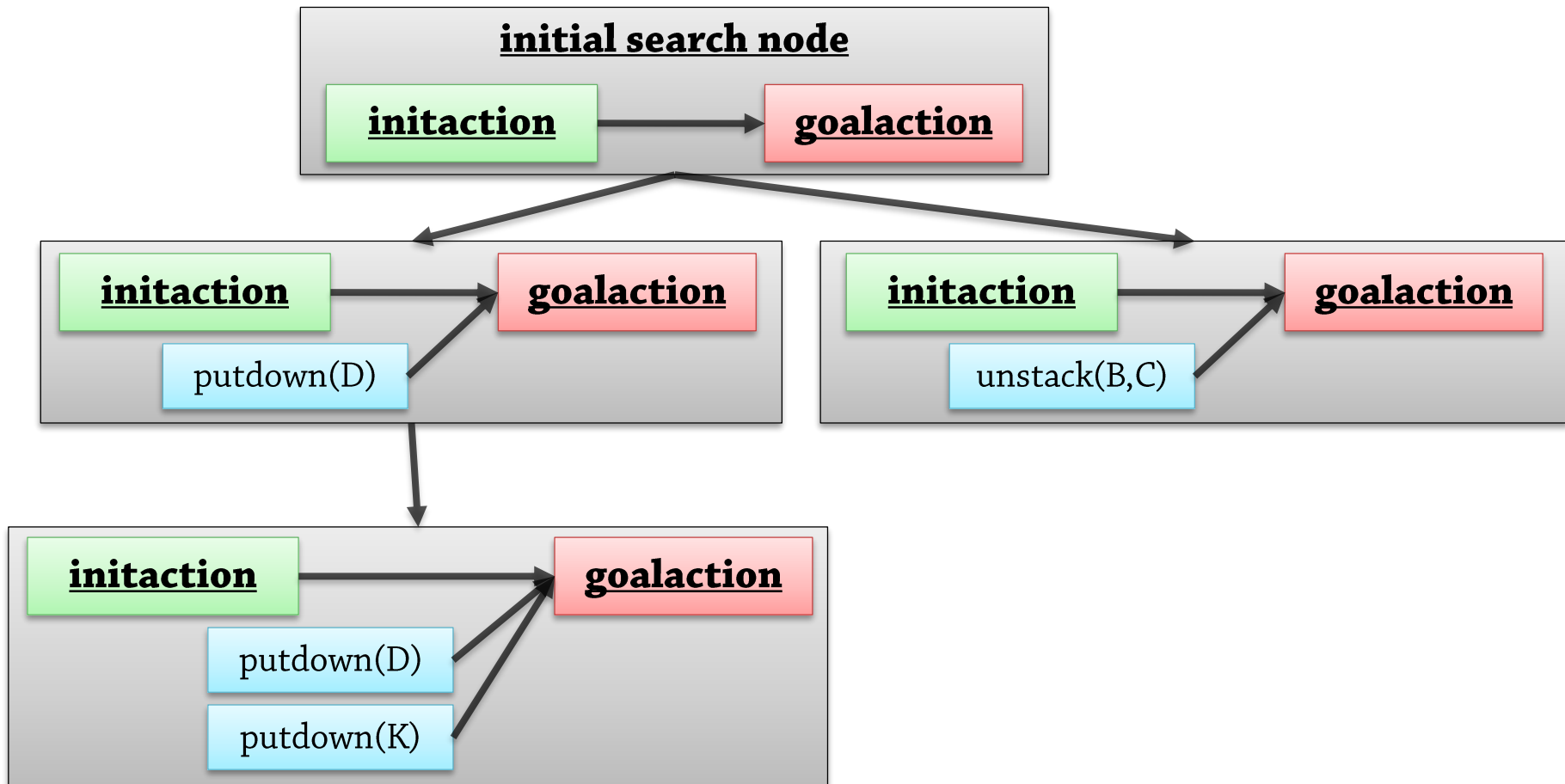
Branching Rule

- We need a **branching rule** as well!

- Forward planning: One successor per action **applicable** in s
- Backward planning: One successor per action **relevant** to g
- POCL planning: One successor for every way that a **flaw in the plan** (**open goal** or **threat**) can be **repaired**



- Gives rise to a search space
 - Use search strategies, backtracking, heuristics, ... to search this space!



**Successors fix Flaws:
Open Goals and Threats**

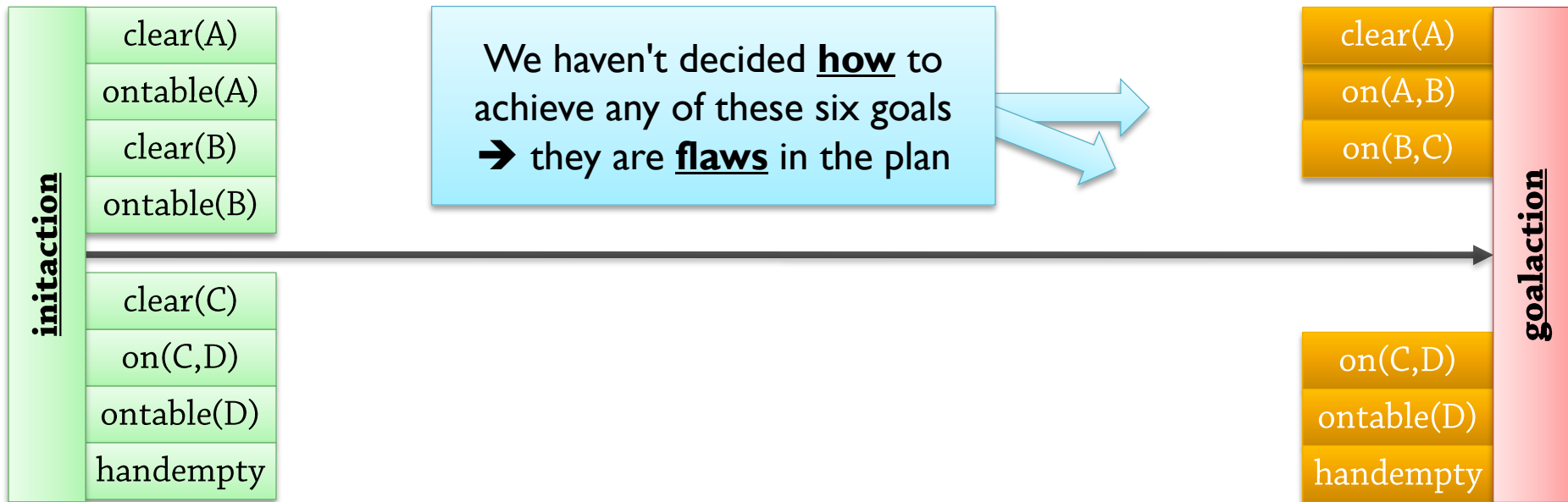
- **Flaw, *noun*:**
 1. a feature that mars the perfection of something; defect; fault: beauty without flaw; the flaws in our plan.
 2. a defect impairing legal soundness or validity.
 3. a crack, break, breach, or rent.

- **Flaw, in POCL planning:**
 - Something we **need to take care of** to **complete the plan**
 - Technical definition: *An open goal or a threat*

- **Not:**
 - Something that has "gone wrong"
 - A problem during planning
 - A mistake in the final solution
 - ...

Flaw Type 1: Open Goals

- Open goal:
 - An action a has a precondition p with no incoming causal link



clear(A) is already true in s_0 , but there is no causal link...

Adding one from s_0 means clear(A) **must never be deleted!**

We need other alternatives too: Delete clear(A), then re-achieve it for goalaction...

Flaw Type 1: Open Goals

- To resolve an open goal :
 - Find an action b that causes p
 - Can be a *new* action
 - Can be an action *already* in the plan, if we can *make* it precede a
 - Add a **causal link**



Partial order! This was not possible in backward search...

Essential:

Even if there is already an action that causes p ,
you can still add a new action that also causes p !

Resolving Open Goals 1

■ Here: **Six** open goals

- Could choose to **find support for clear(A)**:
 - From initaction
 - From a new unstack(B,A), unstack(C,A), or unstack(D,A)
 - From a new stack(A,B), stack(A,C), stack(A,D), or putdown(A)
- Could choose to **find support for on(A,B)**:
 - Only from a new instance of stack(A,B)
- ...

8 distinct
successors

+1 successor

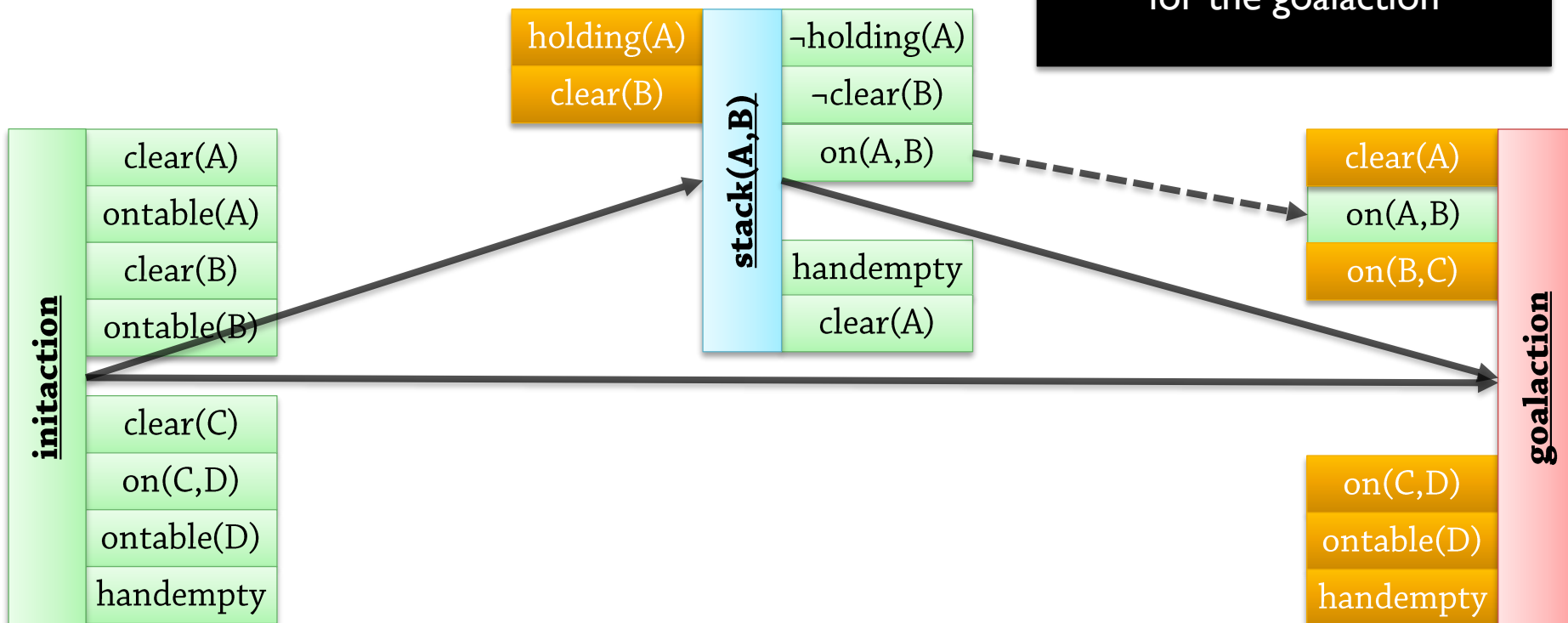


Resolving Open Goals 2

- Suppose we add stack(A,B) to support (achieve) on(A,B)

- Must add a causal link for on(A,B)
 - Dashed line
- Must also add precedence constraints
- Looks totally ordered
 - Because it actually only has one “real” action...

Causal link says:
This instance of stack(A,B)
is responsible for
achieving on(A,B)
for the goalaction



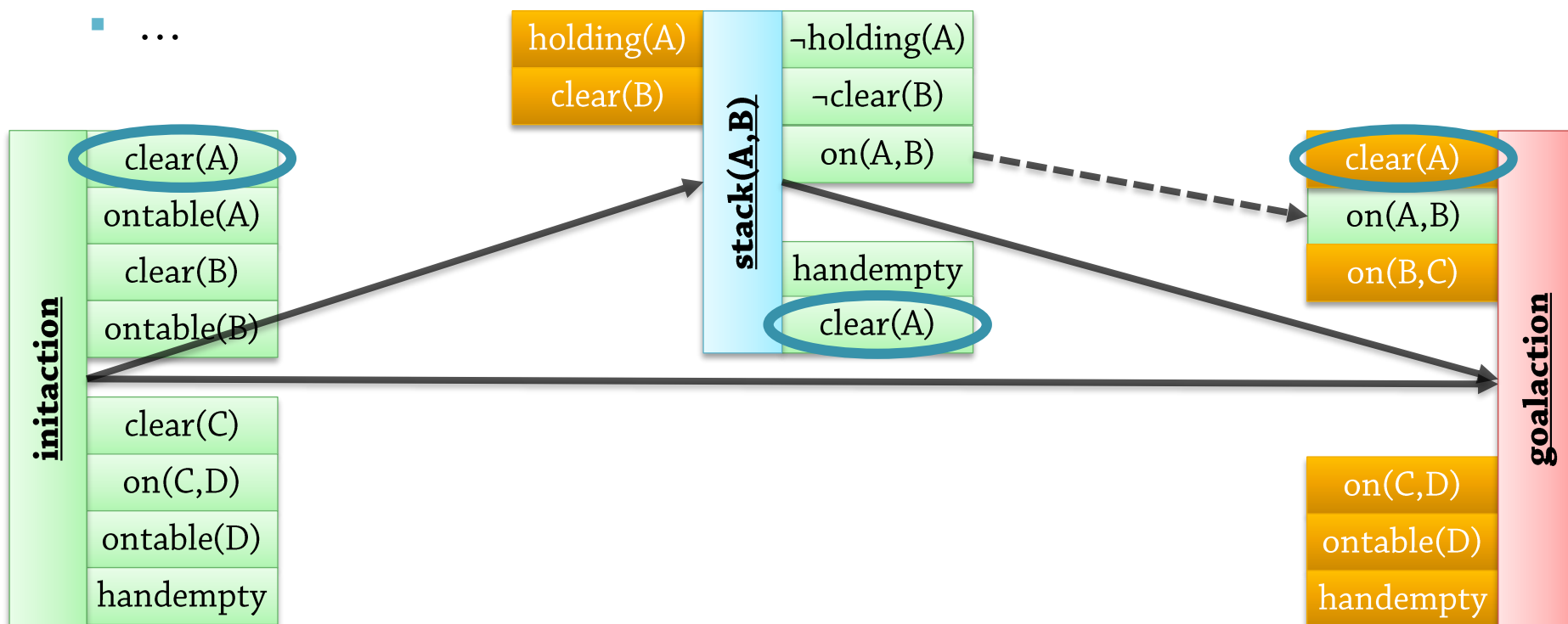
Resolving Open Goals 3

■ Now: **7** open goals (one more!)

■ Can choose to find support for `clear(A)`:

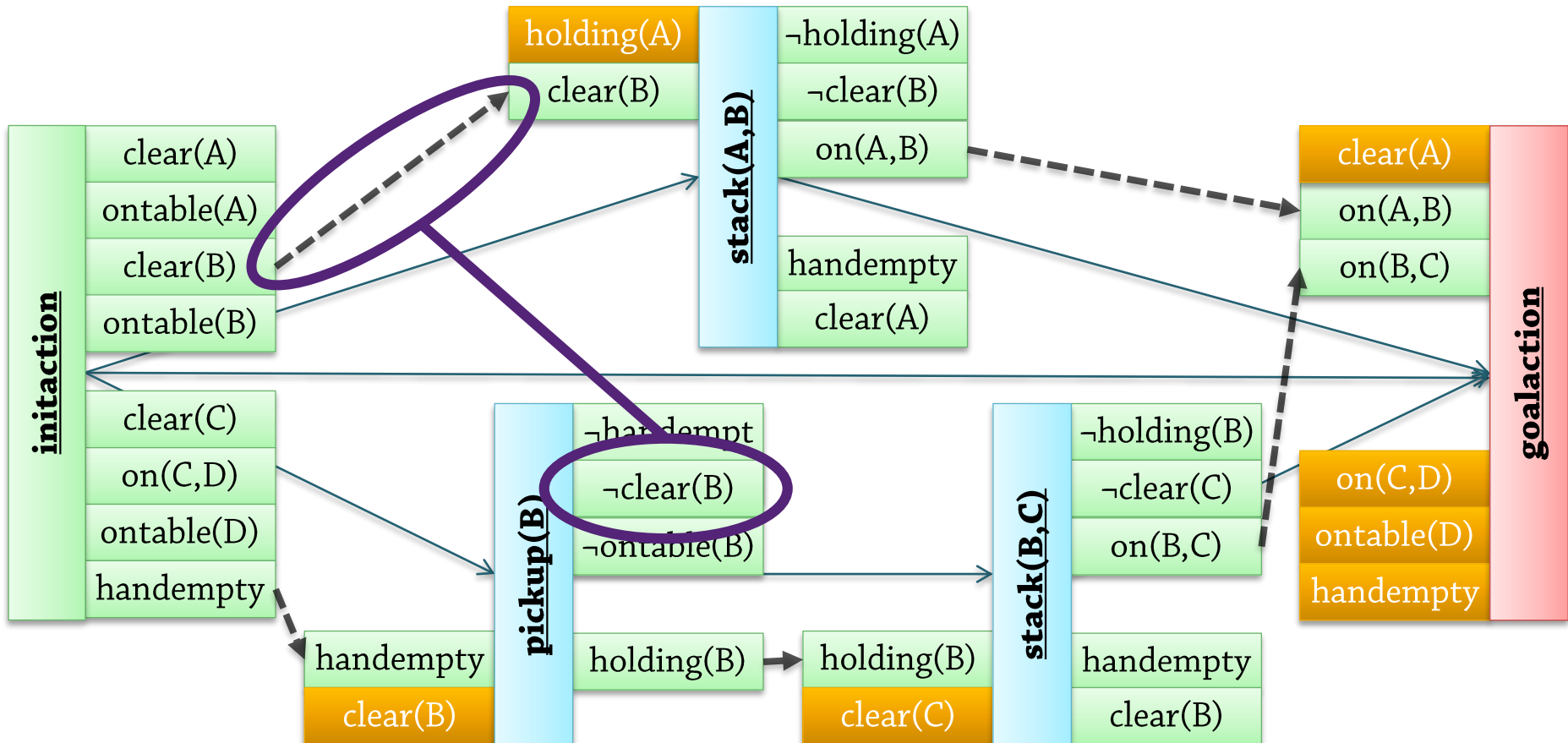
- From the initiation
- From the instance of **stack(A,B)** that we just added
- From a **new** instance of **stack(A,B)**, `stack(A,C)`, `stack(A,D)`, or `putdown(A)`
- From a **new** instance of `unstack(B,A)`, `unstack(C,A)`, or `unstack(D,A)`

■ ...



- Second flaw type: A threat

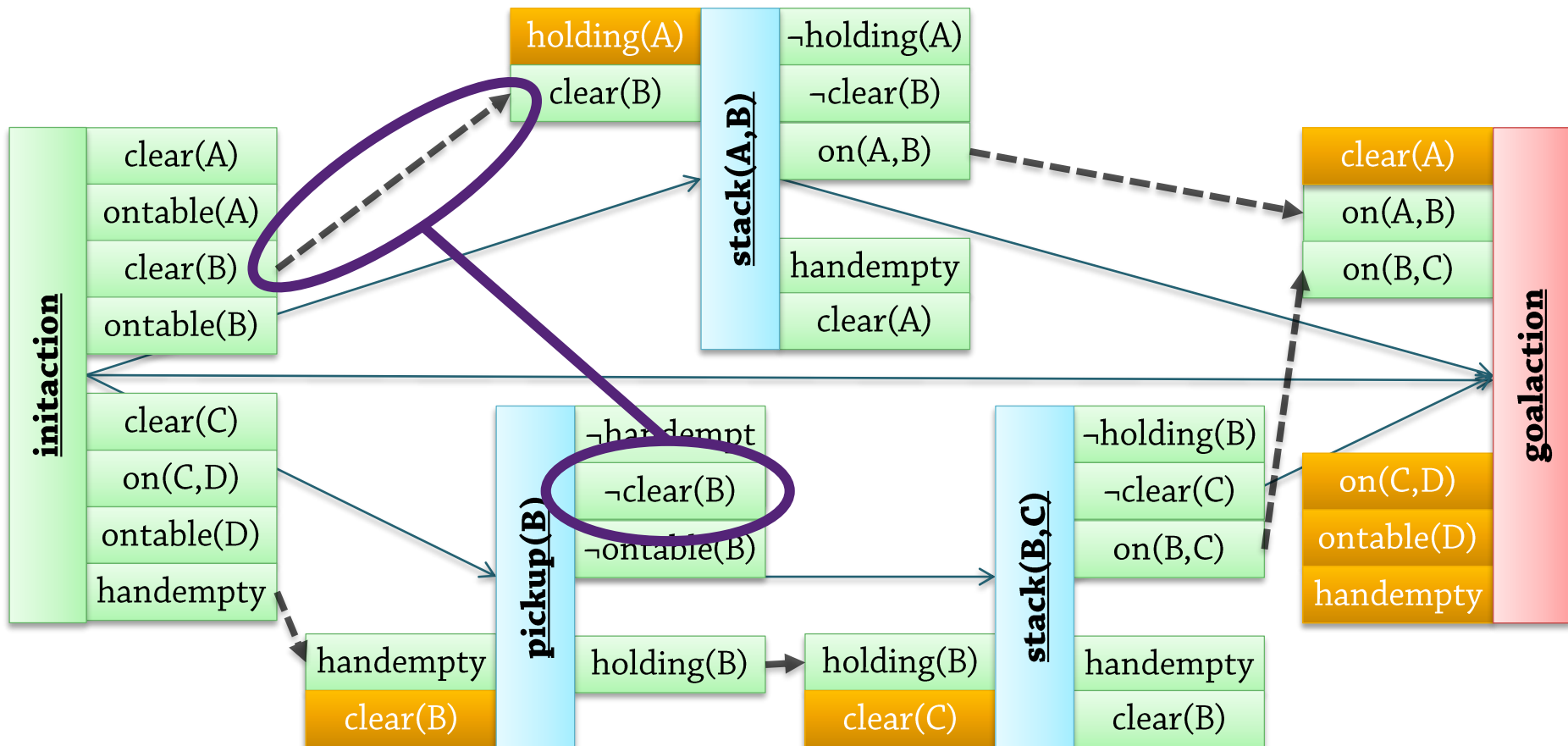
- initiation **supports** clear(B) for stack(A,B) – there's a causal link
- pickup(B) **deletes** clear(B), and may occur **between** initiation and stack(A,B)
- So we can't be certain that clear(B) still holds when stack(A,B) starts!



Flaw Type 2: Threats (2)

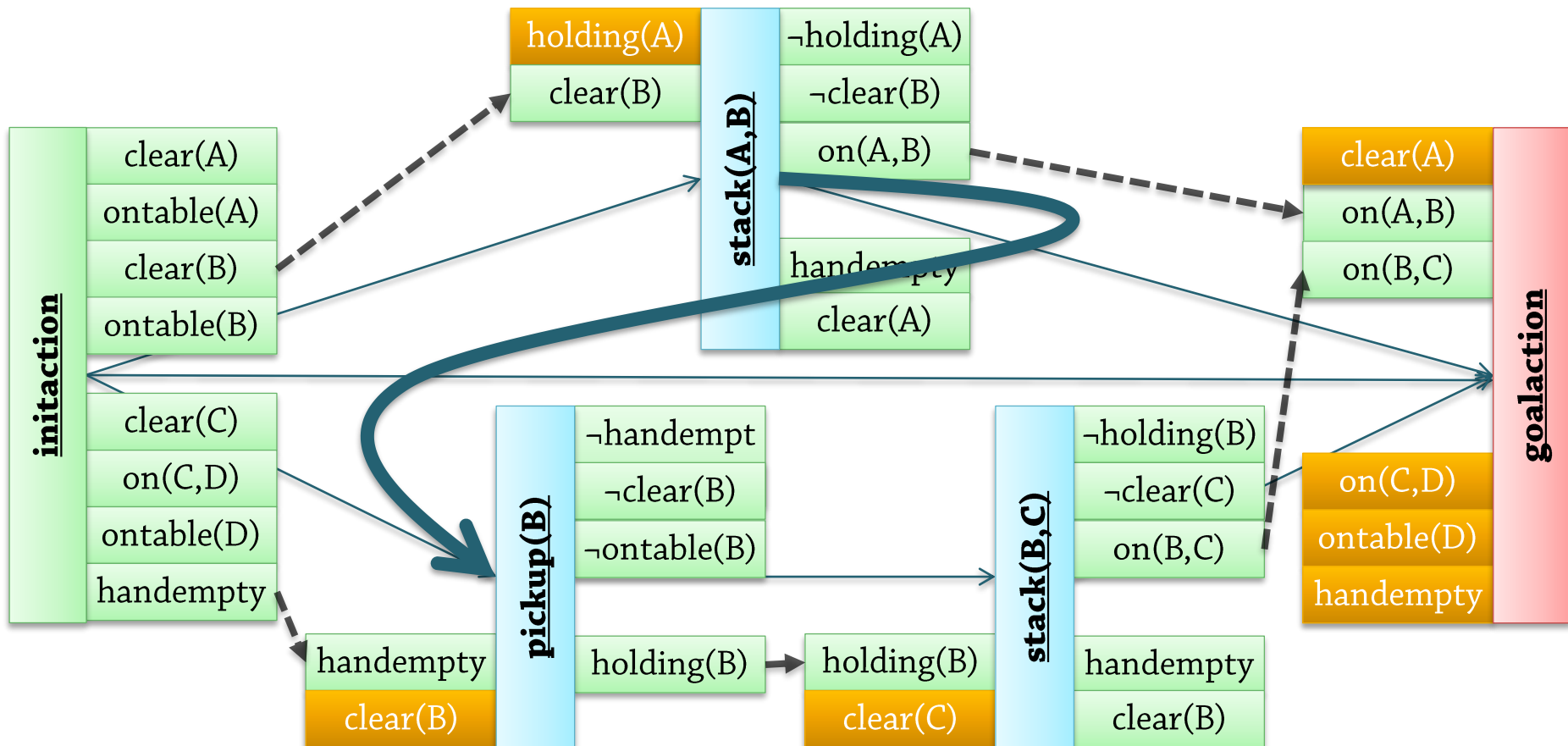
- Some possible execution orders:

- <..., stack(A,B), pickup(B), ...> -- preconditions of stack(A,B) OK
- <..., pickup(B), stack(A,B), ...> -- preconditions of stack(A,B) not satisfied



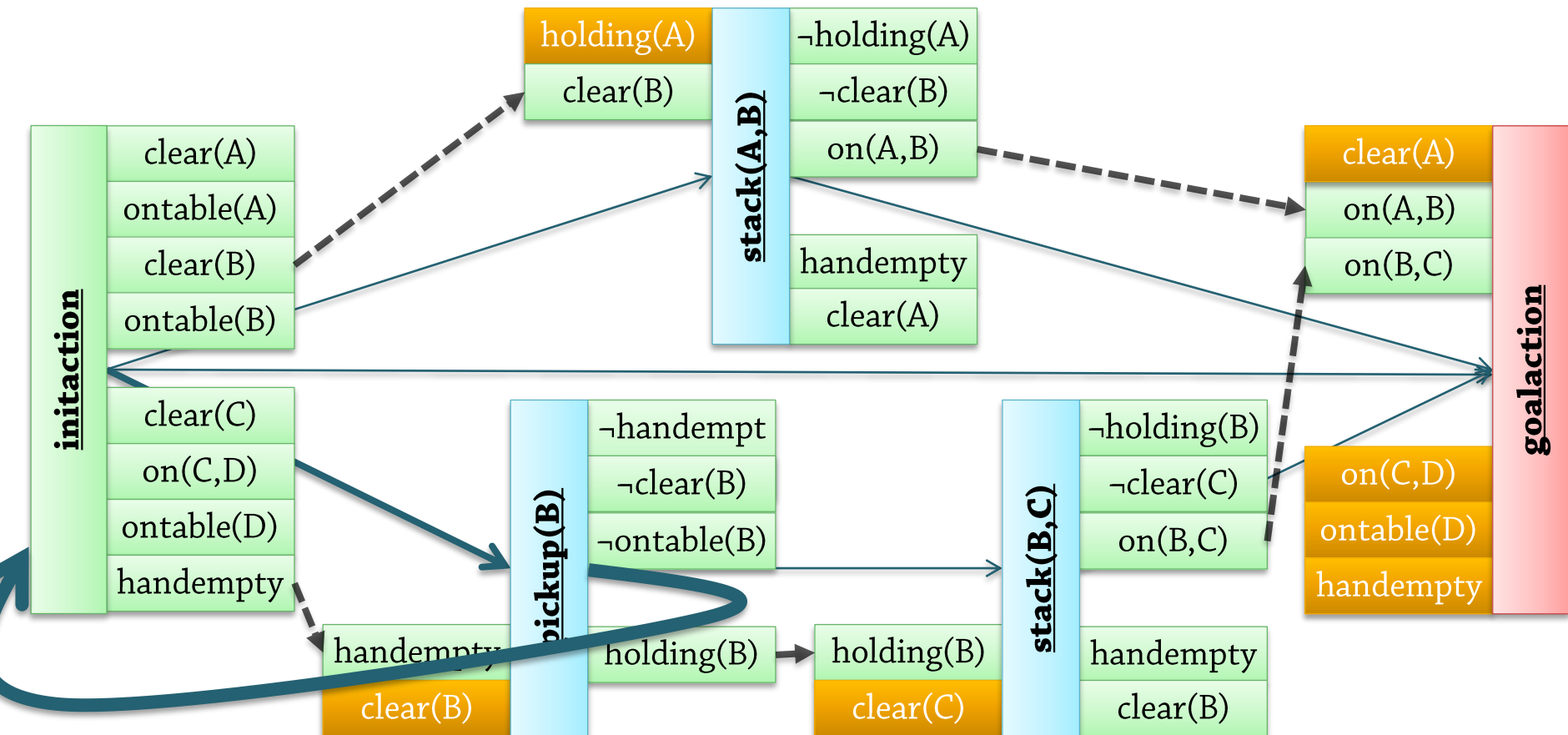
Resolving Threats 1

- How to **make sure** that $\text{clear}(B)$ holds when $\text{stack}(A,B)$ starts?
 - Alternative 1: The action that **disturbs** the precondition is placed **after** the action that **has** the precondition
 - Only possible if the resulting partial order is consistent (acyclic)!



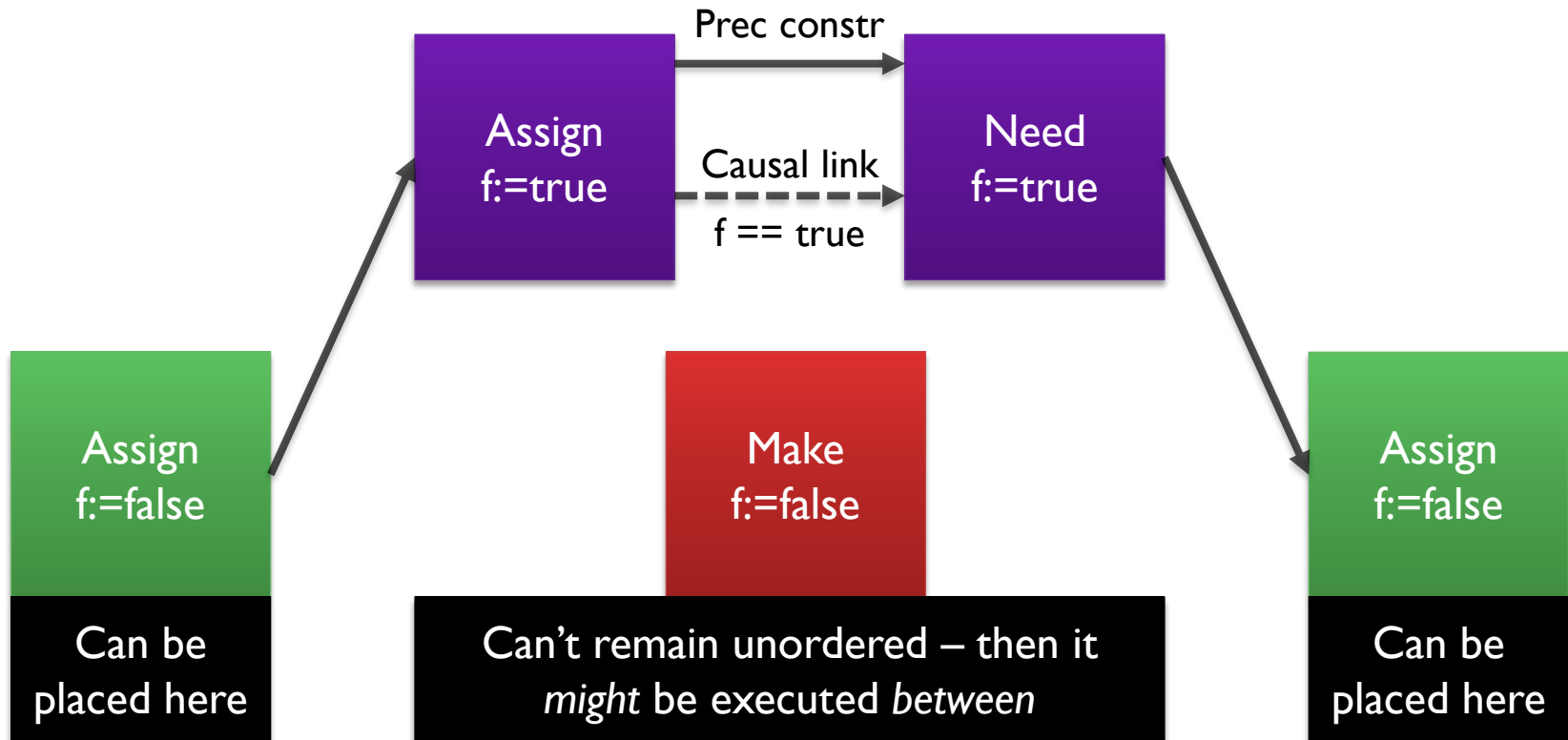
Resolving Threats 2

- Alternative 2:
 - The action that **disturbs** the precondition is placed **before** the action that **supports** the precondition
 - Only possible if the resulting partial order is consistent – not in this case!



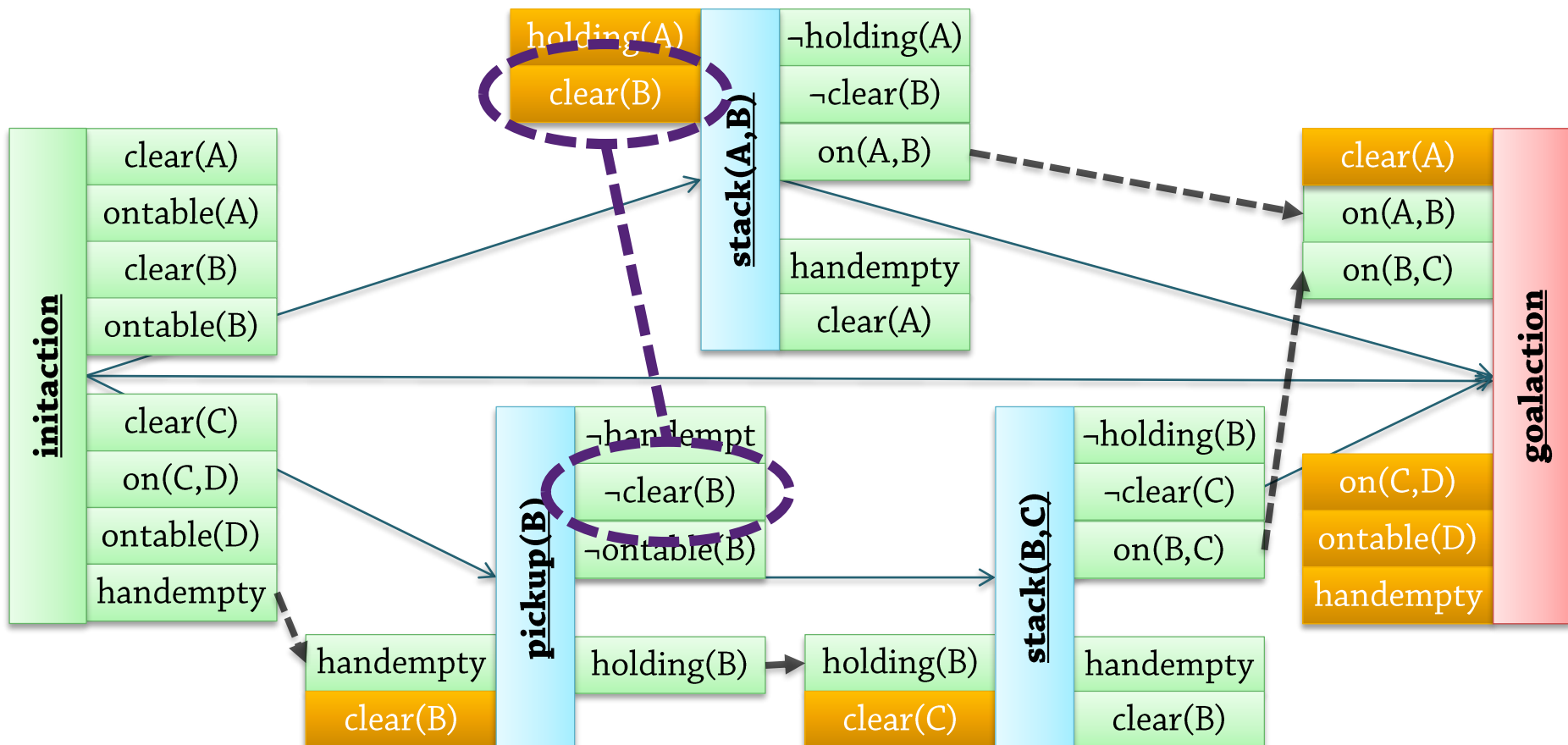
Resolving Threats 3

- Summary:



Resolving Threats 4

- Only **causal links** can be threatened!
 - Below, pickup(B) does **not** threaten the precondition clear(B) of stack(A,B)
 - We haven't decided yet **how** to achieve clear(B): No incoming causal link
 - So we can't claim that its achievement is threatened!



POCL Algorithms

■ POCL planning – one possible formulation (sound/complete):

■ planner(initial, goal):

open $\leftarrow$ { Make-initial-plan(initial, goal) }

while open $\neq \emptyset$:

select a plan π in open

open $\leftarrow$ open – { π }

flaws $\leftarrow$ OpenGoals(π) $\cup$ Threats(π)

if flaws = $\emptyset$ **then**

return π

else

select any flaw $\varphi \in$ flaws

resolvers $\leftarrow$ FindResolvers(φ , π) // May be the empty set!

foreach r in resolvers:

$\pi' \leftarrow$ Refine(r, π) // Actually apply the resolver

open $\leftarrow$ open $\cup$ { π' }

endwhile

return failure

Requires
heuristics

**Can prove: The plan is correct
iff there are no remaining flaws
(no open goals, no threats)**

Requires
heuristics

One flaw chosen – no backtracking!

But all resolvers must be tested...

Returns a **partially** ordered solution plan
Any **total** ordering of this plan will achieve the goals

- **Plan-Space Planning (book)**: Same principle, slightly different

- PSP(π)

$flaws \leftarrow \text{OpenGoals}(\pi) \cup \text{Threats}(\pi)$

if $flaws = \emptyset$ **then return** π

select any flaw $\varphi \in flaws$

$resolvers \leftarrow \text{FindResolvers}(\varphi, \pi)$

if $resolvers = \emptyset$ **then return** failure // Backtrack

nondeterministically choose a resolver $\rho \in resolvers$

$\pi' \leftarrow \text{Refine}(\rho, \pi)$ // Actually apply the resolver

return PSP(π')

end

Nondeterministic choice
instead of "foreach resolver",
list of open nodes

- Call PSP(the initial plan)

Understanding Backtracking in POCL

```
planner(initial, goal):  
  open  $\leftarrow$  { Make-initial-plan(initial, goal) }  
  while open  $\neq \emptyset$ :  
    select a plan  $\pi$  in open  
    open  $\leftarrow$  open - {  $\pi$  }  
    flaws  $\leftarrow$  OpenGoals( $\pi$ )  $\cup$  Threats( $\pi$ )  
    if flaws =  $\emptyset$  then  
      return  $\pi$   
    else  
      select any flaw  $\varphi \in$  flaws  
      resolvers  $\leftarrow$  FindResolvers( $\varphi$ ,  $\pi$ )  
      foreach r in resolvers:  
         $\pi' \leftarrow$  Refine( $\rho$ ,  $\pi$ ) // Actually apply resolver  
        open  $\leftarrow$  open  $\cup$  {  $\pi'$  }  
  endwhile  
  return failure
```

If we choose a random flaw to resolve and then can't find a plan, we don't have to try another flaw. **Why?**

- 1) Every flaw *has* to be resolved
- 2) Choosing this flaw *later* cannot help us resolve it
- 3) Choosing this flaw *later* cannot help us resolve some other flaw

We must be able to try different resolvers.
Why?

Choosing one resolver *can* prevent other problem resolutions.

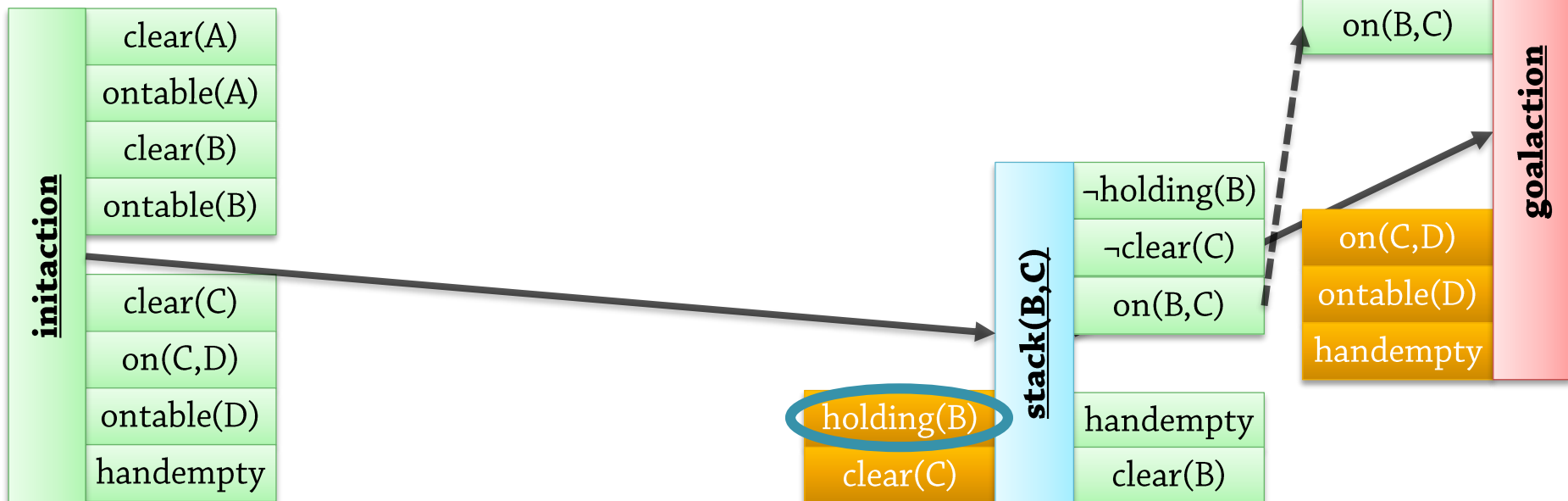
- Open goal: Use action A or B?
- Threat: Which order to choose?

Lifted Planning: Partial Action Instantiation

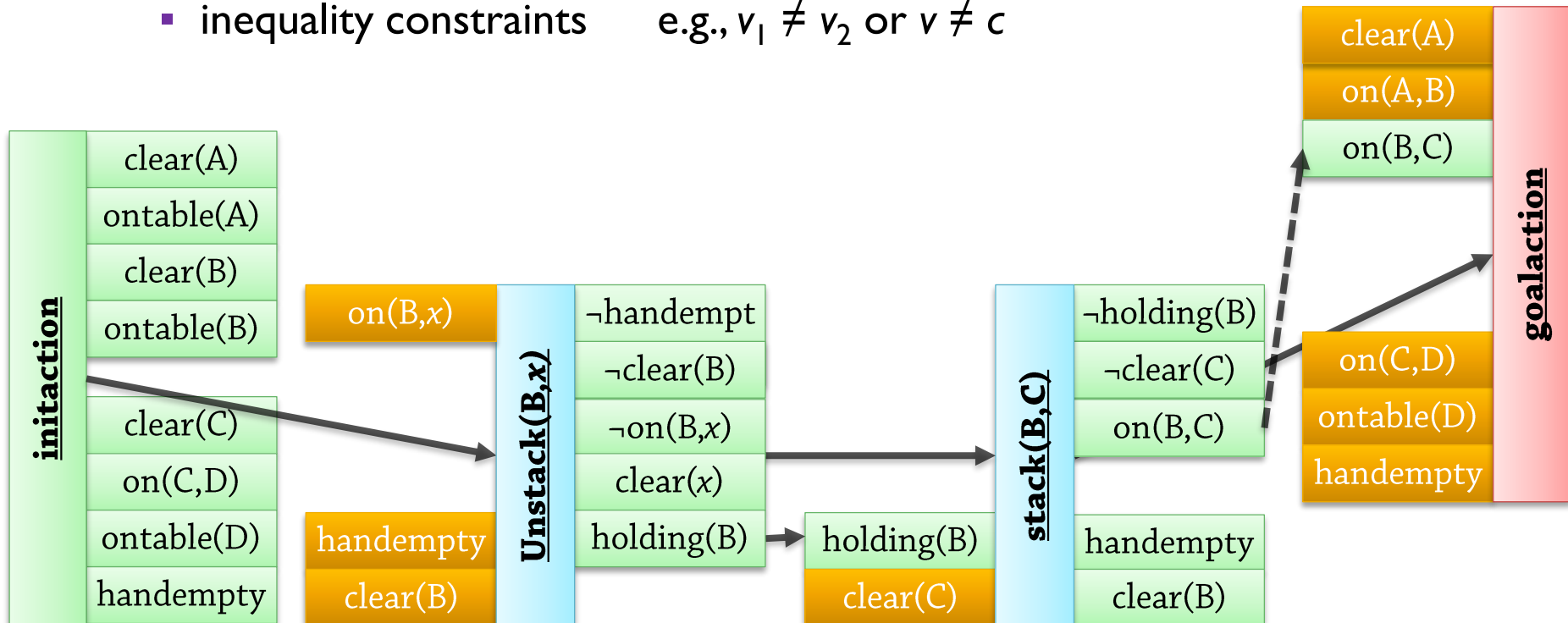
Partial Instantiation

- Suppose we want to achieve `holding(B)`
 - **Ground** search generates **many** alternatives
 - Add `unstack(B,A)`, `unstack(B,F)`, `unstack(B,G)`, ...
 - Add `pickup(B)`
 - Let's take the idea of **least commitment** one step further
 - **Lifted** search generates two **partially instantiated** alternatives
 - Add **unstack(B, x)**
 - Add `pickup(B)`

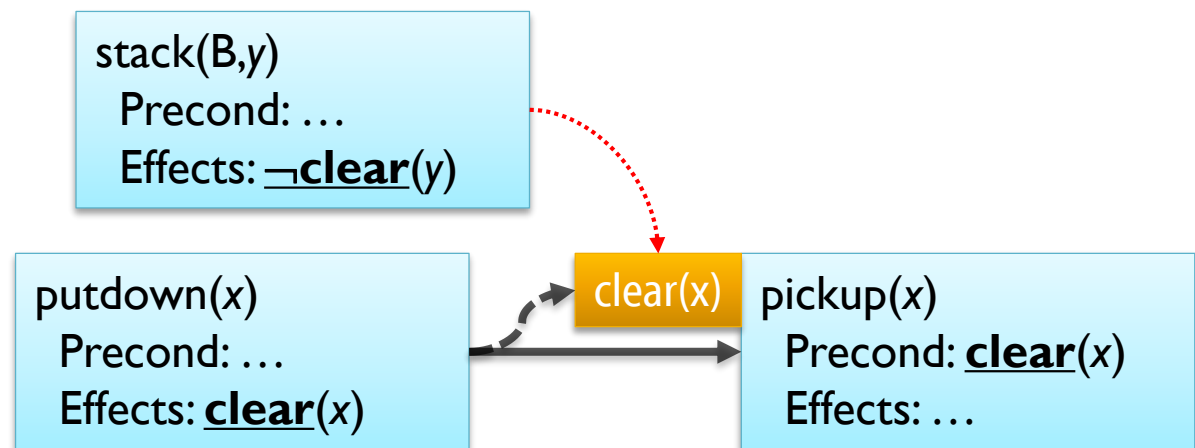
So far, we see no reason why we should unstack B from any **specific** block!



- A **lifted partial-order plan** consists of:
 - A set of **possibly unground actions**
 - A set of **precedence constraints**: a must precede b
 - A set of **causal links**: action a establishes the precondition p needed by b
 - A set of **binding constraints**:
 - equality constraints e.g., $v_1 = v_2$ or $v = c$
 - inequality constraints e.g., $v_1 \neq v_2$ or $v \neq c$



- Another way of resolving threats for lifted plans:
 - For partly uninstantiated actions, we may find potential threats
 - $\text{stack}(B,y)$ may threaten the causal link, but only if $x=y$
 - Can be resolved by adding a constraint: $x \neq y$

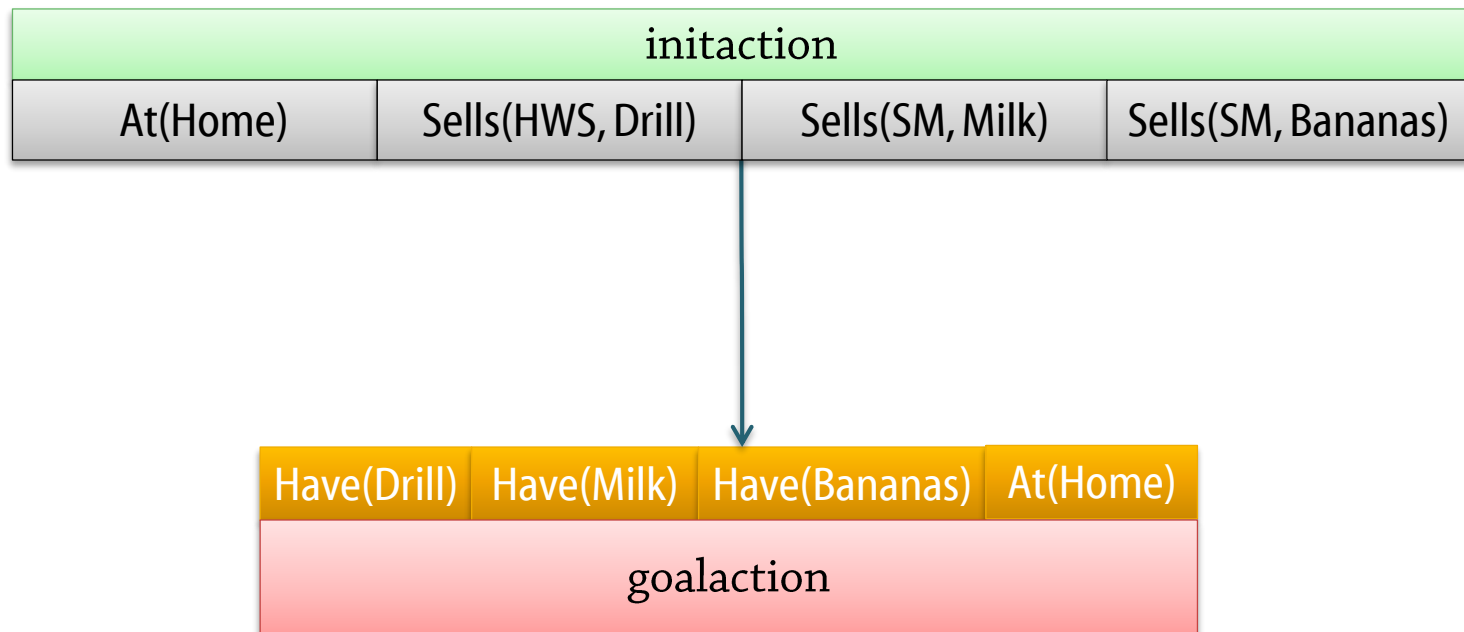


Complete Example

- Running Example: Similar to an example in ALMA
 - Russell and Norvig's *Artificial Intelligence: A Modern Approach* (1st ed.)
 - Operator **Go(from,to)**
 - Precond: $\text{At}(\text{from})$
 - Effects: $\text{At}(\text{to}), \neg\text{At}(\text{from})$
 - Operator **Buy(product, store)**
 - Precond: $\text{At}(\text{store}), \text{Sells}(\text{store}, \text{product})$
 - Effects: $\text{Have}(\text{product})$
 - **Initial state**
 - $\text{At}(\text{Home}), \text{Sells}(\text{HWS}, \text{Drill}), \text{Sells}(\text{SM}, \text{Milk}), \text{Sells}(\text{SM}, \text{Bananas})$
 - **Goal**
 - $\text{At}(\text{Home}), \text{Have}(\text{Drill}), \text{Have}(\text{Milk}), \text{Have}(\text{Bananas})$

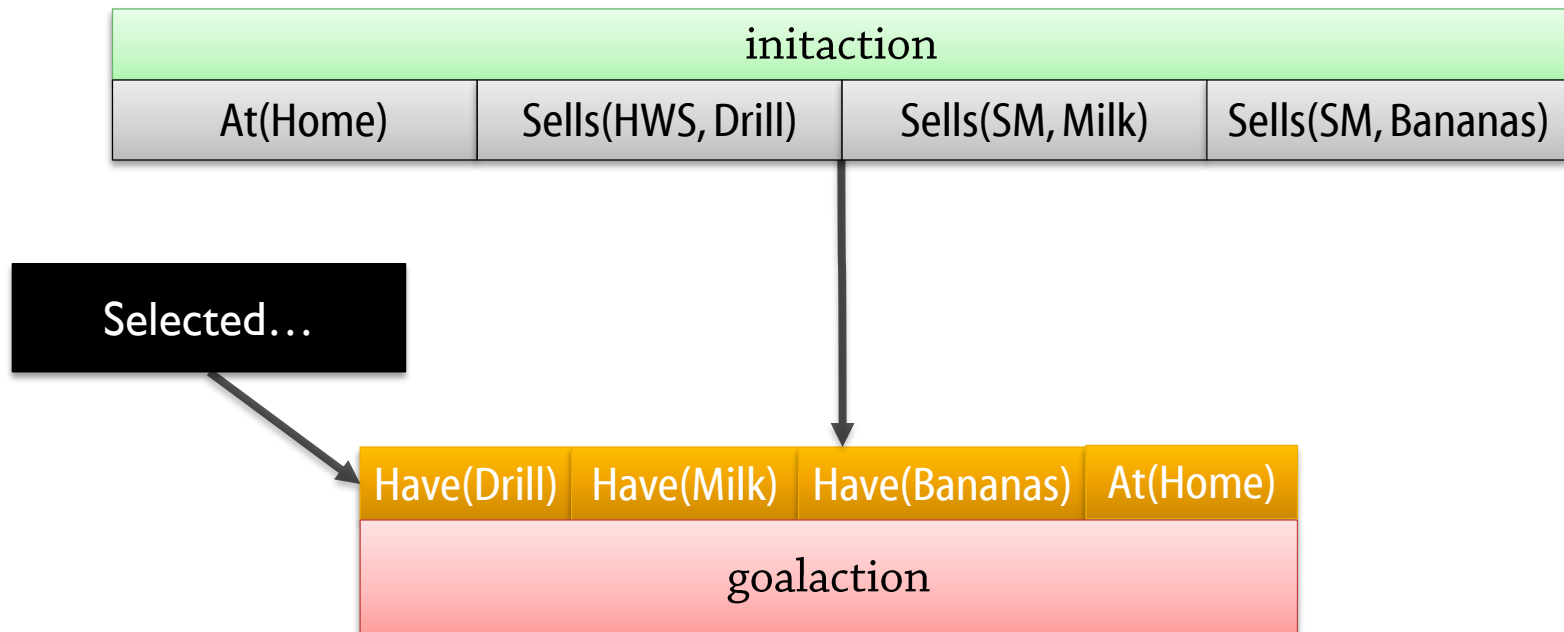
Example (continued)

- PSP takes a plan Π as its argument
 - Initial plan: **initaction**, **goalaction**, and an ordering constraint



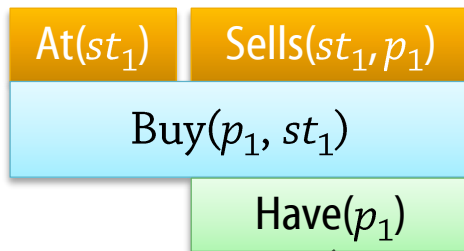
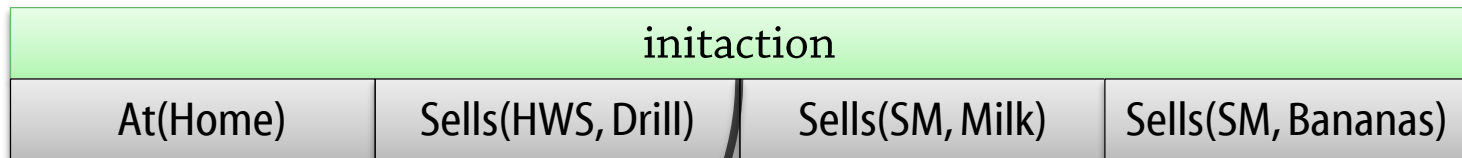
Example (continued)

- Four **flaws** exist: Open goals
 - Suppose our heuristics tell us to resolve Have(Drill) first



Example (continued)

- Have(drill) is not achieved by any action in the current plan
- But **Buy(product, store)** achieves Have(product)
 - Partially instantiate:
Buy(Drill, store) (right now we don't care where we buy it)



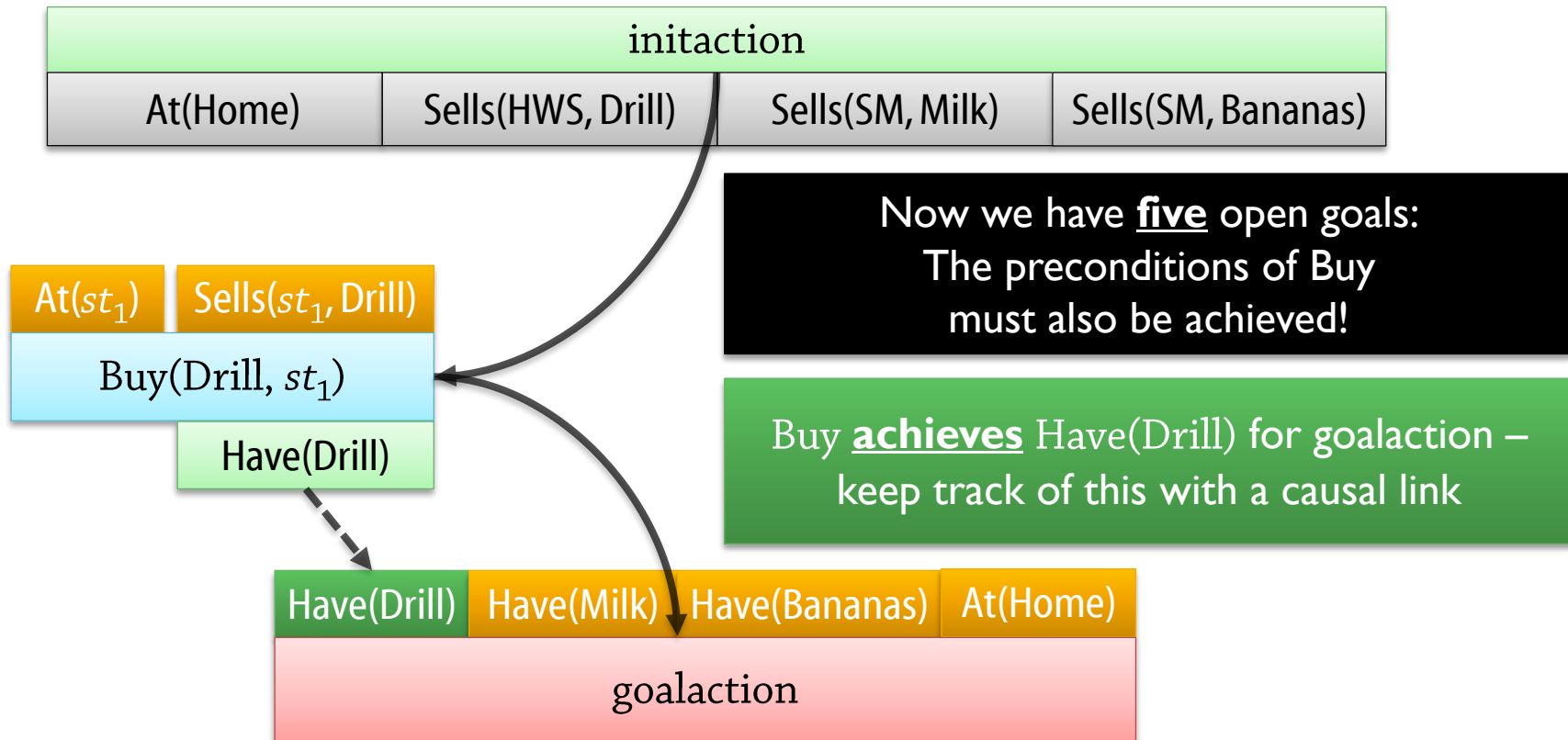
Buy achieves Have(Drill) for goalaction – keep track of this with a causal link

Bindings:
 $p_1 = \text{Drill}$



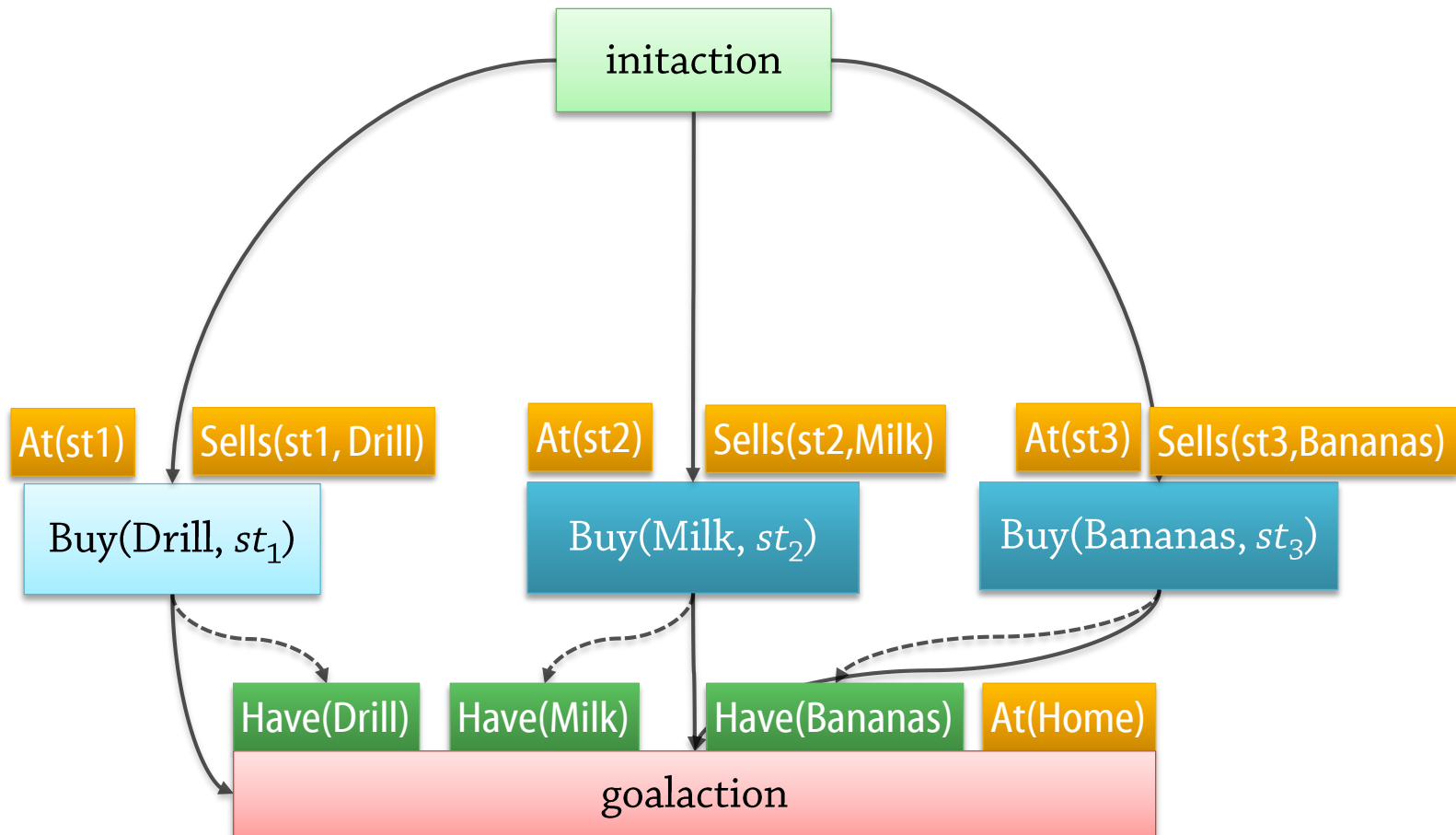
Example (continued)

- **Alternative Notation** for simplicity
 - Variable bindings are implicit in the diagram



Example (continued)

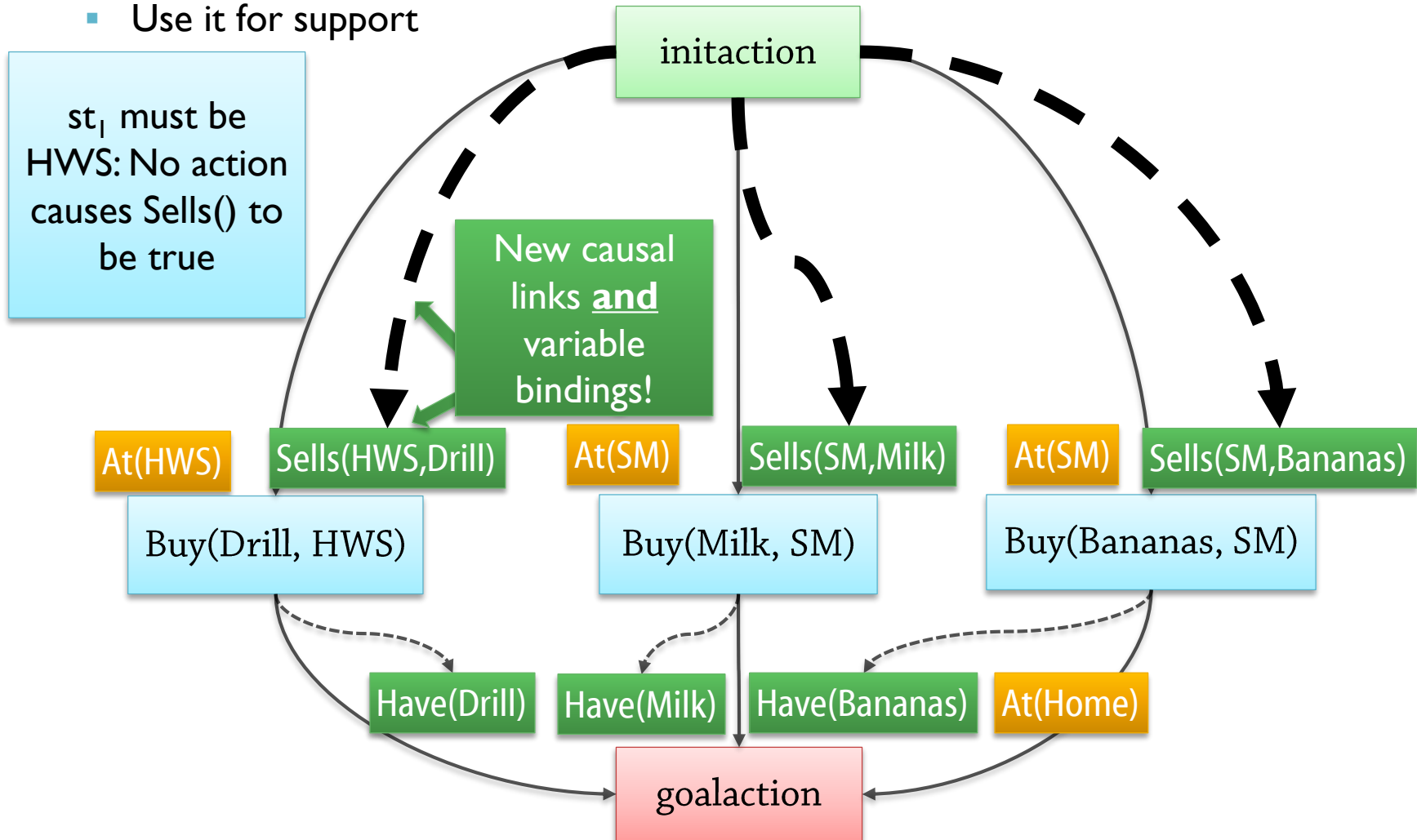
- The first **three** refinement steps
 - These are the only possible ways to establish the Have preconditions
 - We don't care in which **order** we buy things!



Example (continued)

- Three more refinement steps

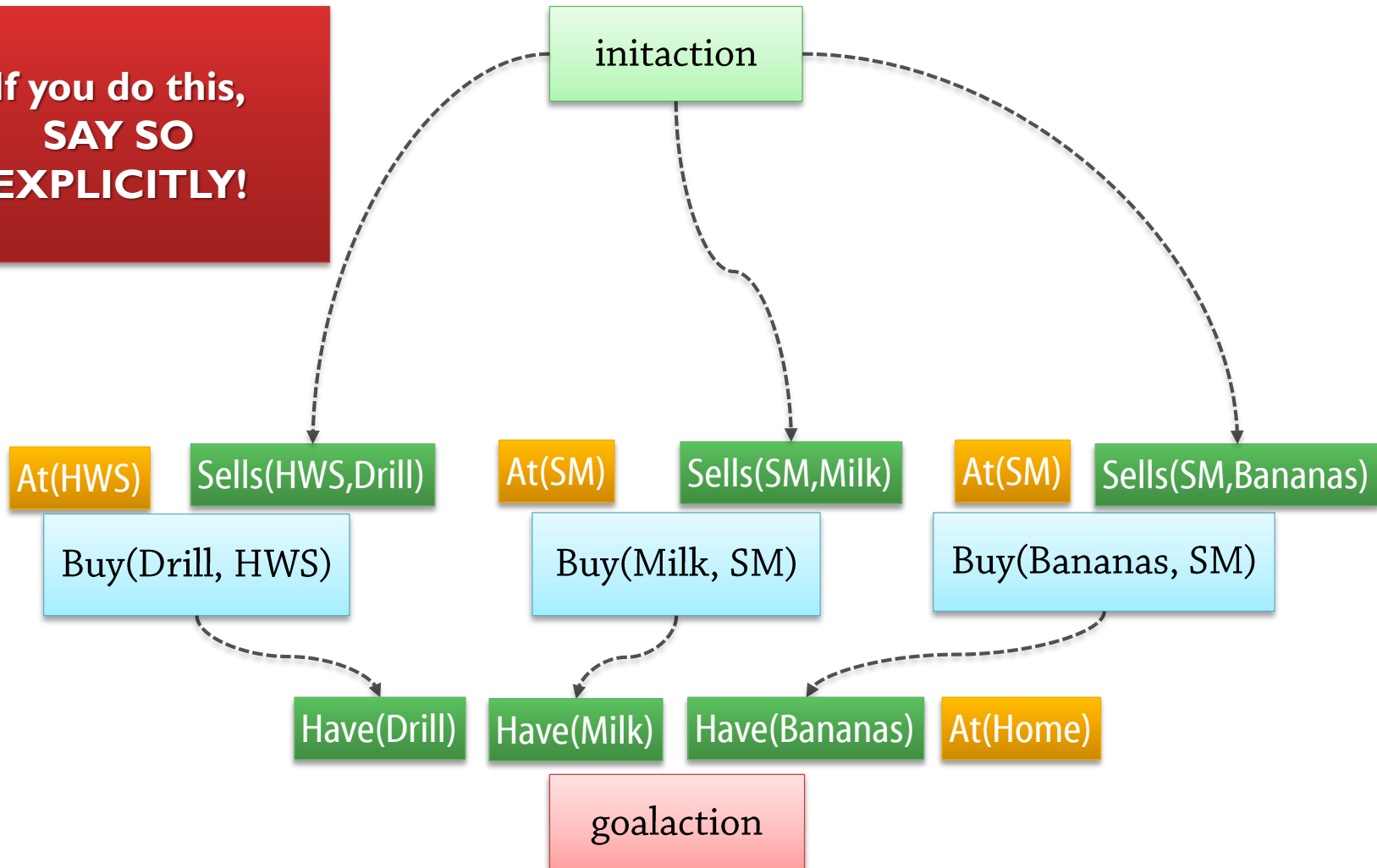
- No action causes Sells(...) to be true – except the “fake” initial action!
- Use it for support



Example (continued)

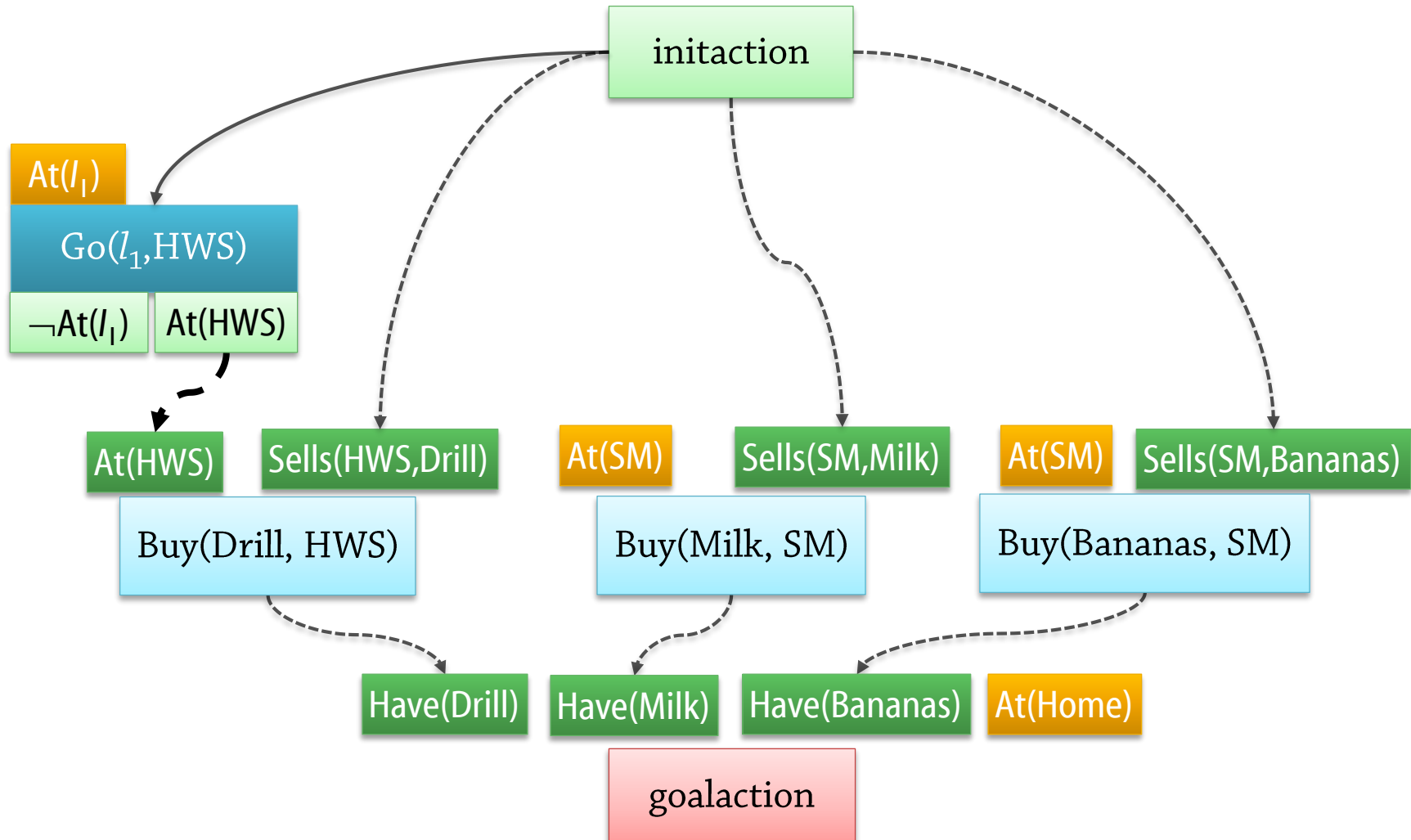
- It's getting messy!
 - Let's omit the precedence constraints that are implicit in causal links...

If you do this,
**SAY SO
EXPLICITLY!**



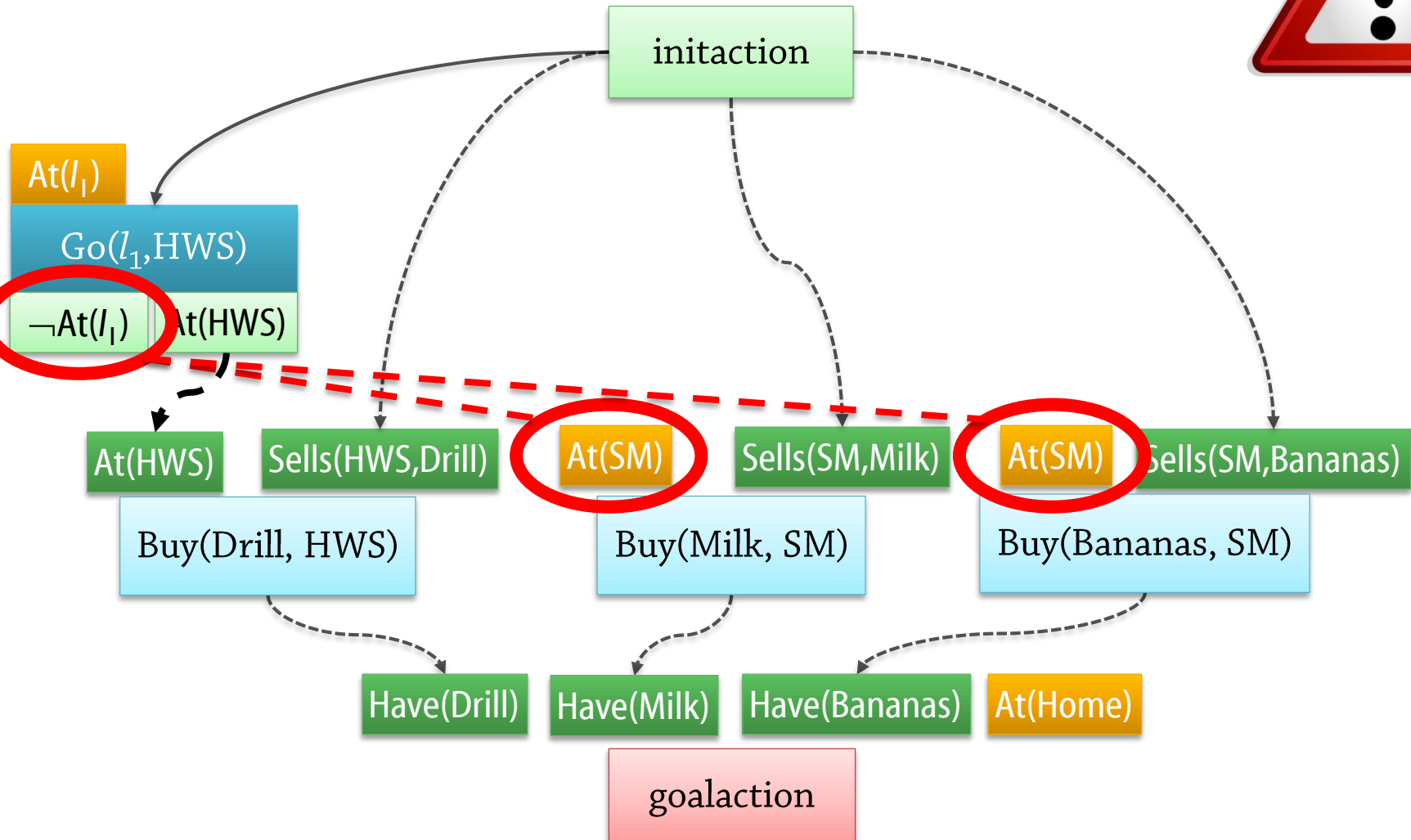
Example (continued)

- To establish $At(HWS)$: **Must** go there from somewhere



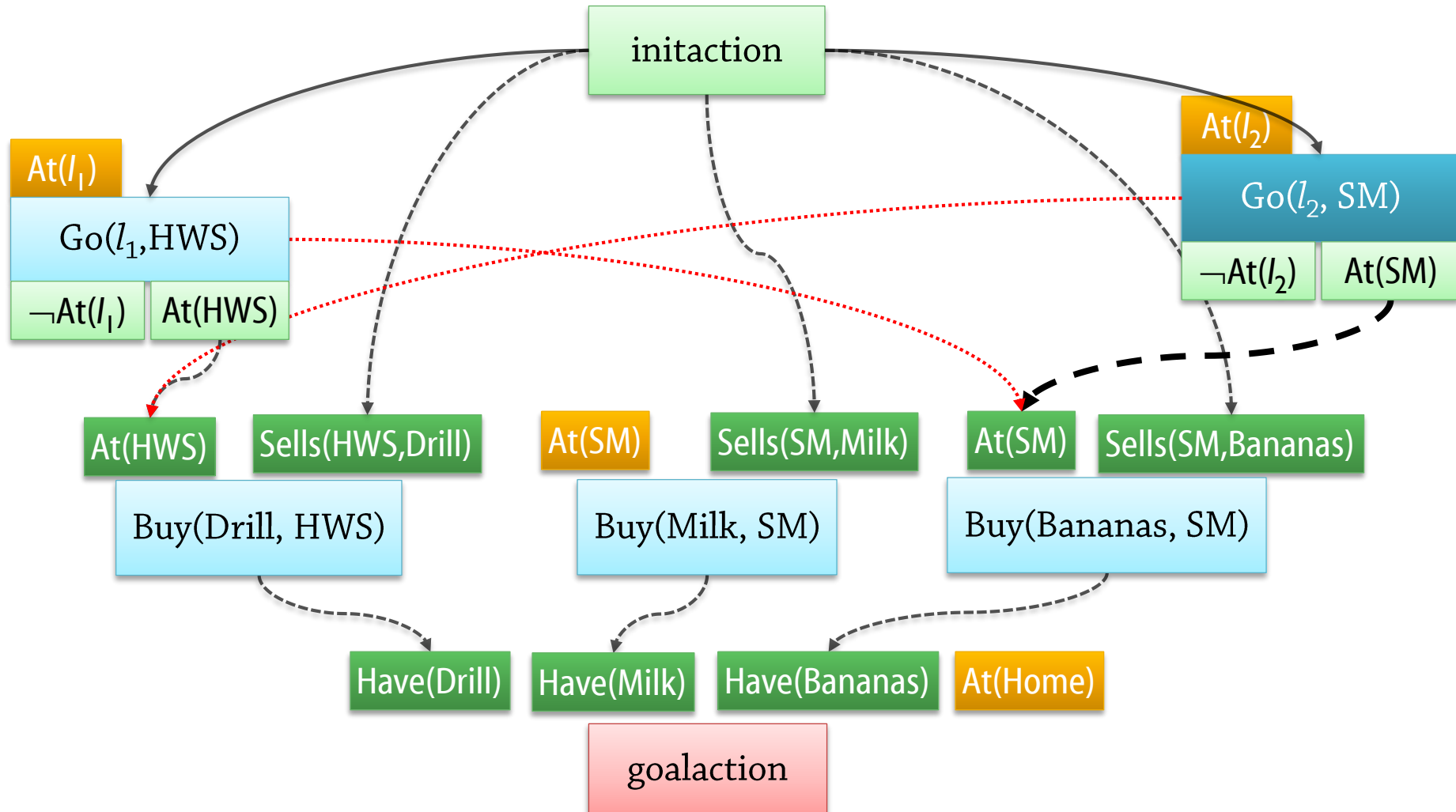
Example (continued)

- Does $\neg \text{at}(I_1)$ threaten $\text{At}(\text{SM})$?
 - No! Only a **causal link** to $\text{At}(\text{SM})$ can be threatened



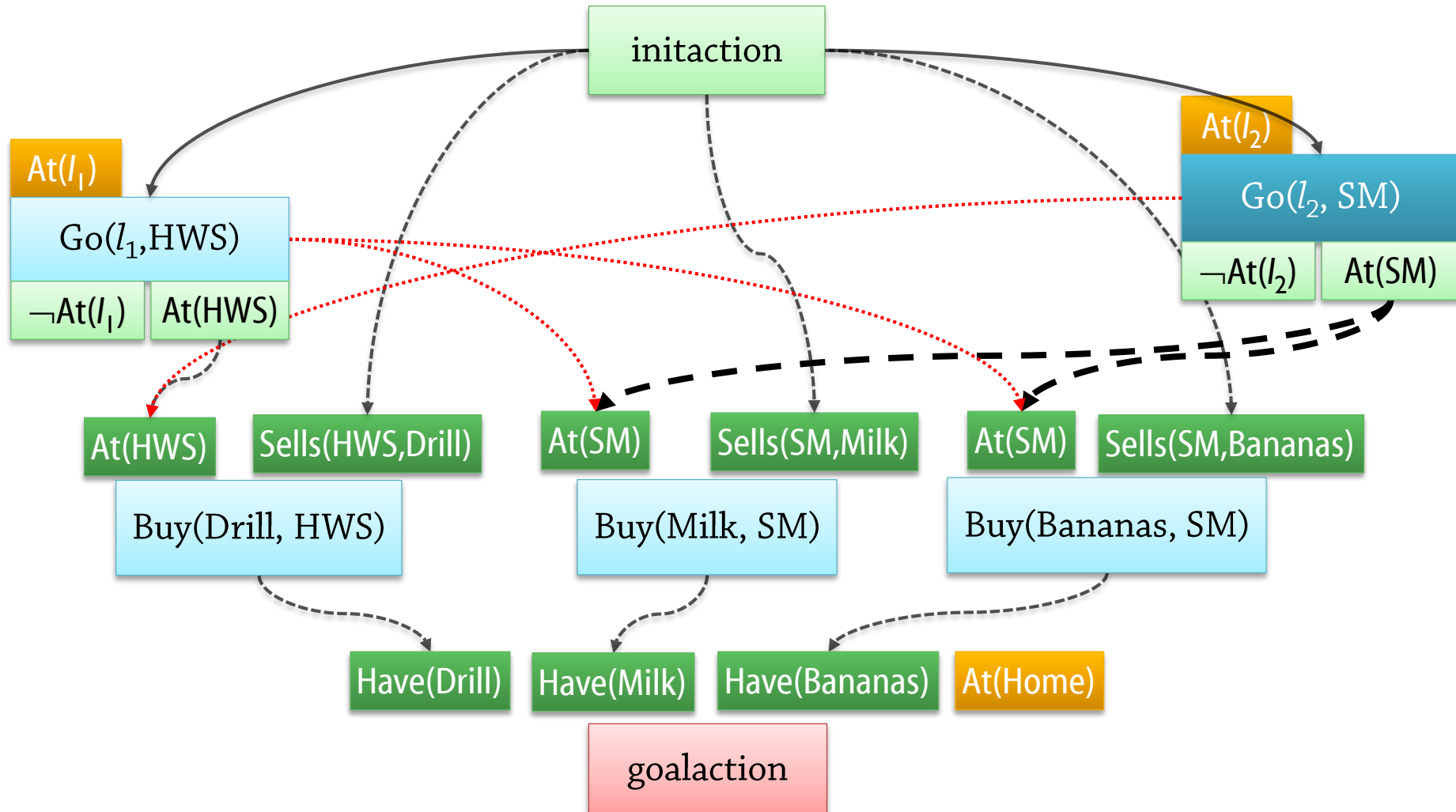
Example (continued)

- To establish $At(SM)$: Must go there from somewhere
 - Mutual threats...



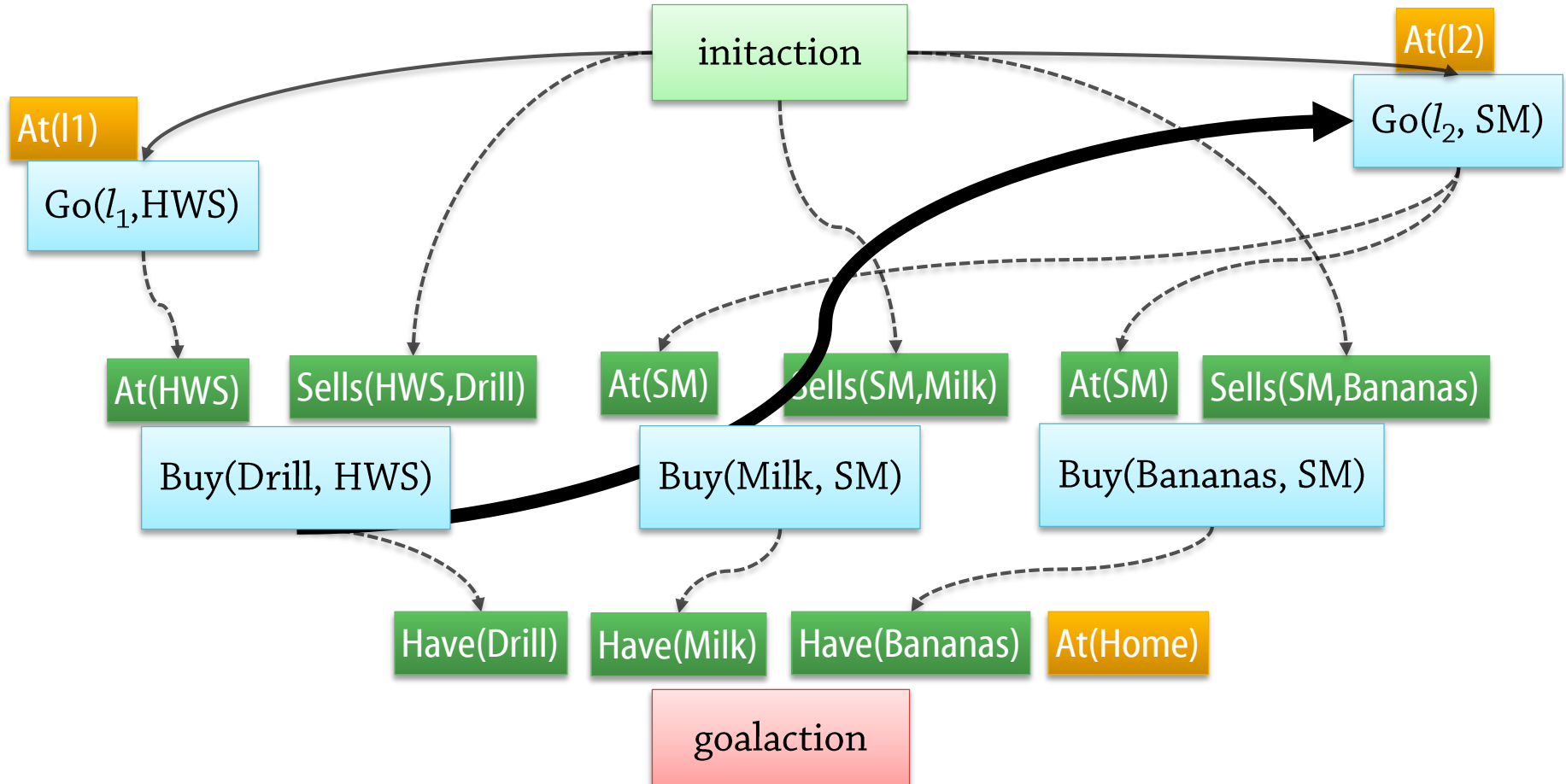
Example (continued)

- Let's use the same action for **both** $\text{At}(\text{SM})$ preconditions...
 - More threats – could deal with them now or wait



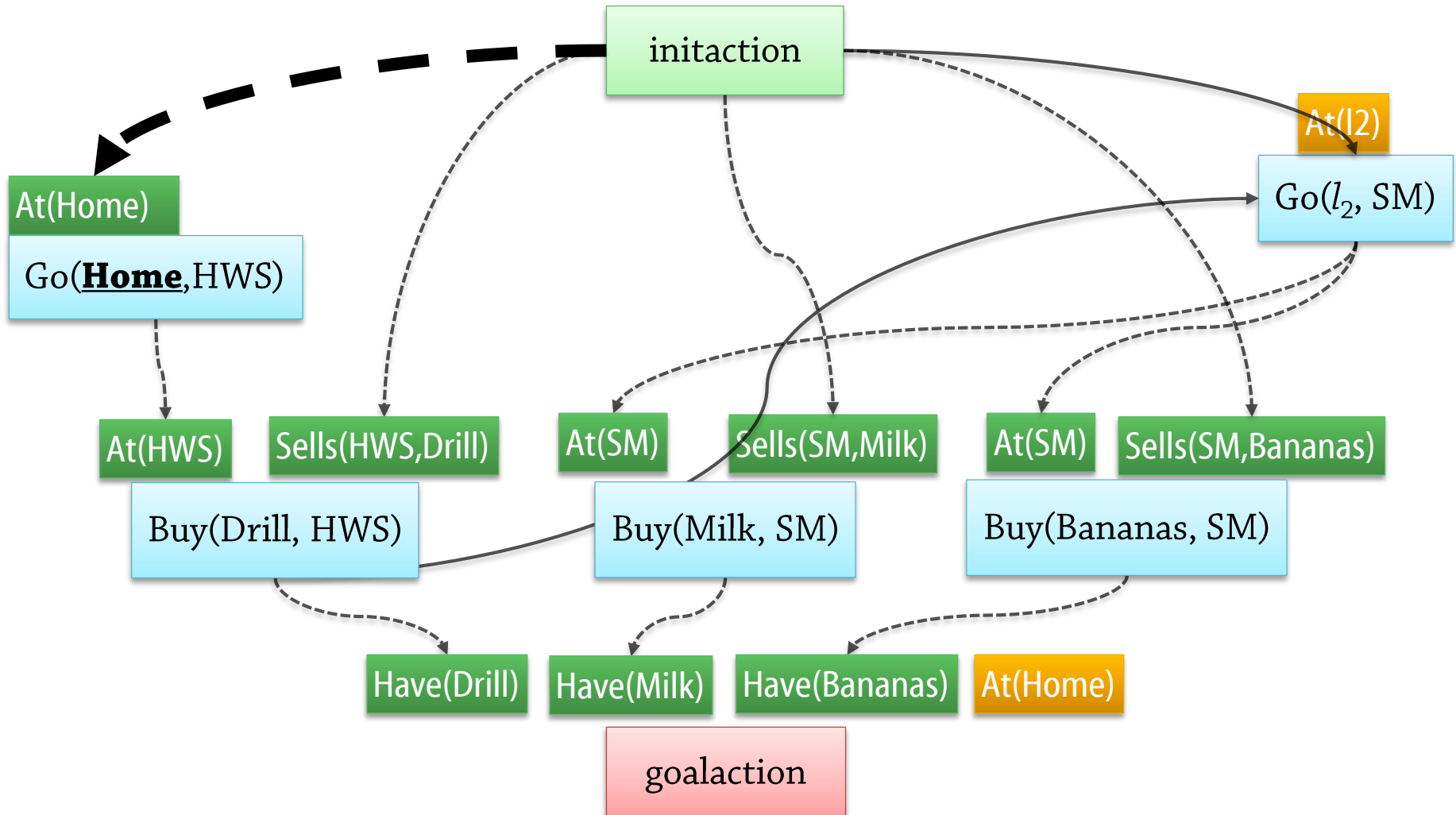
Example (continued)

- **Nondet. choice**: how to resolve the threat to $At(HWS)$?
 - Our choice: make the “requirer” $Buy(Drill)$ precede the “threatener” $Go(l_2, SM)$
 - Also happens to resolve the other two threats
 - “Threatener” $Go(l_1, HWS)$ before “achiever” $Go(l_2, SM)$



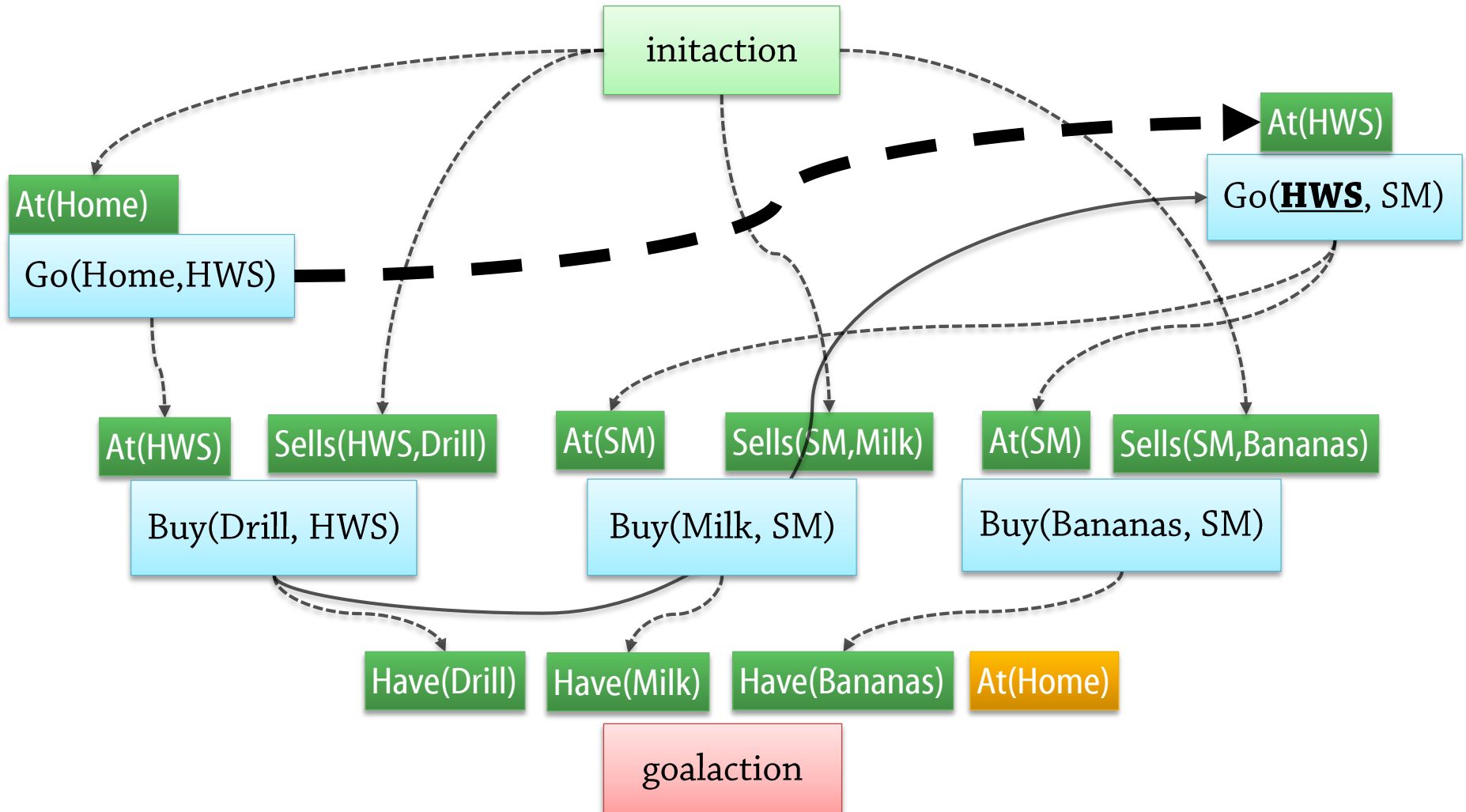
Example (continued)

- **Nondet. choice**: how to establish $At(l_1)$?
 - We'll do it from initaction, with $l_1 = \text{Home}$



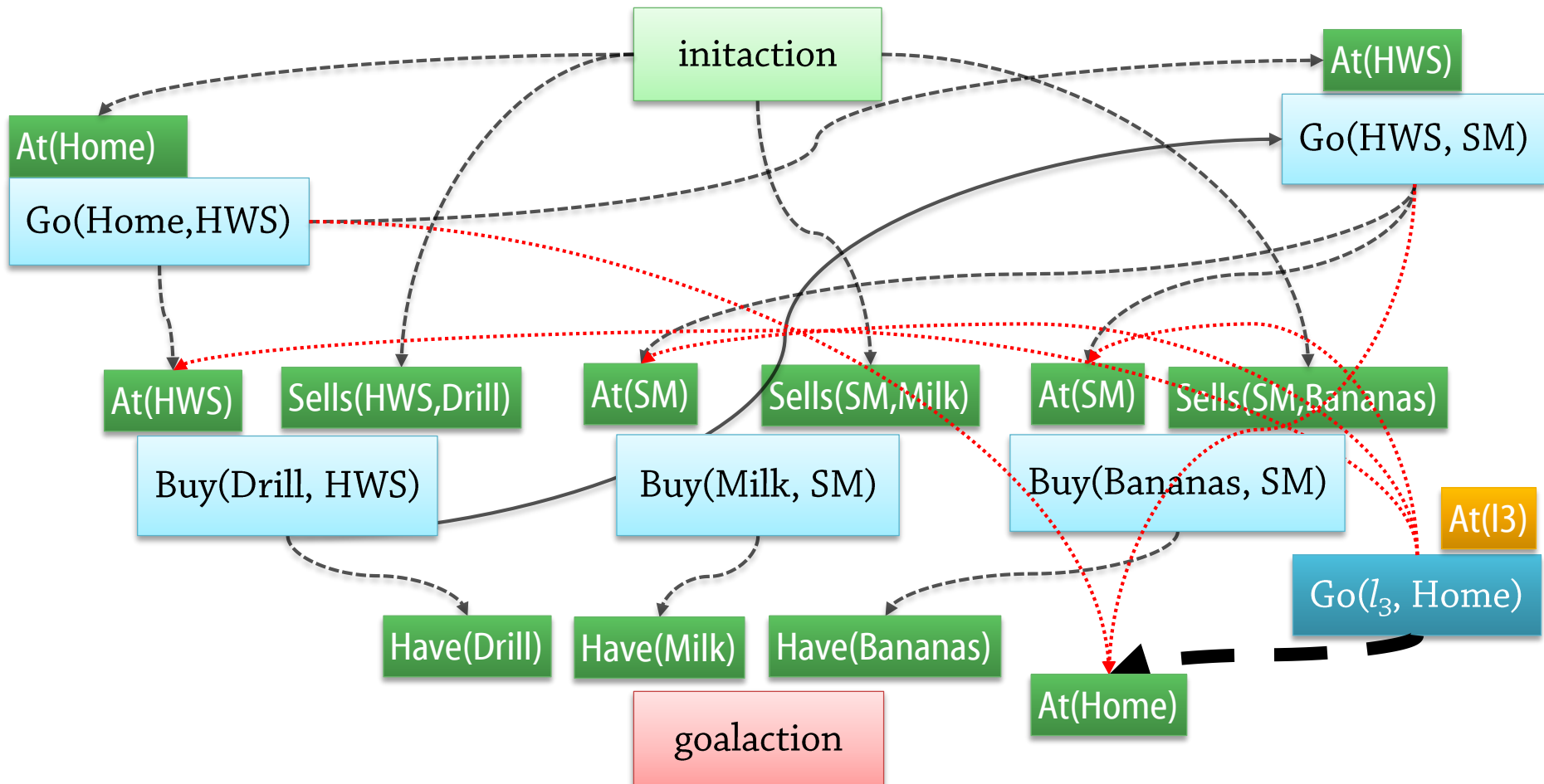
Example (continued)

- **Nondeterministic choice**: how to establish $At(l_2)$?
 - We'll do it from $Go(Home, HWS)$, with $l_2 = HWS$



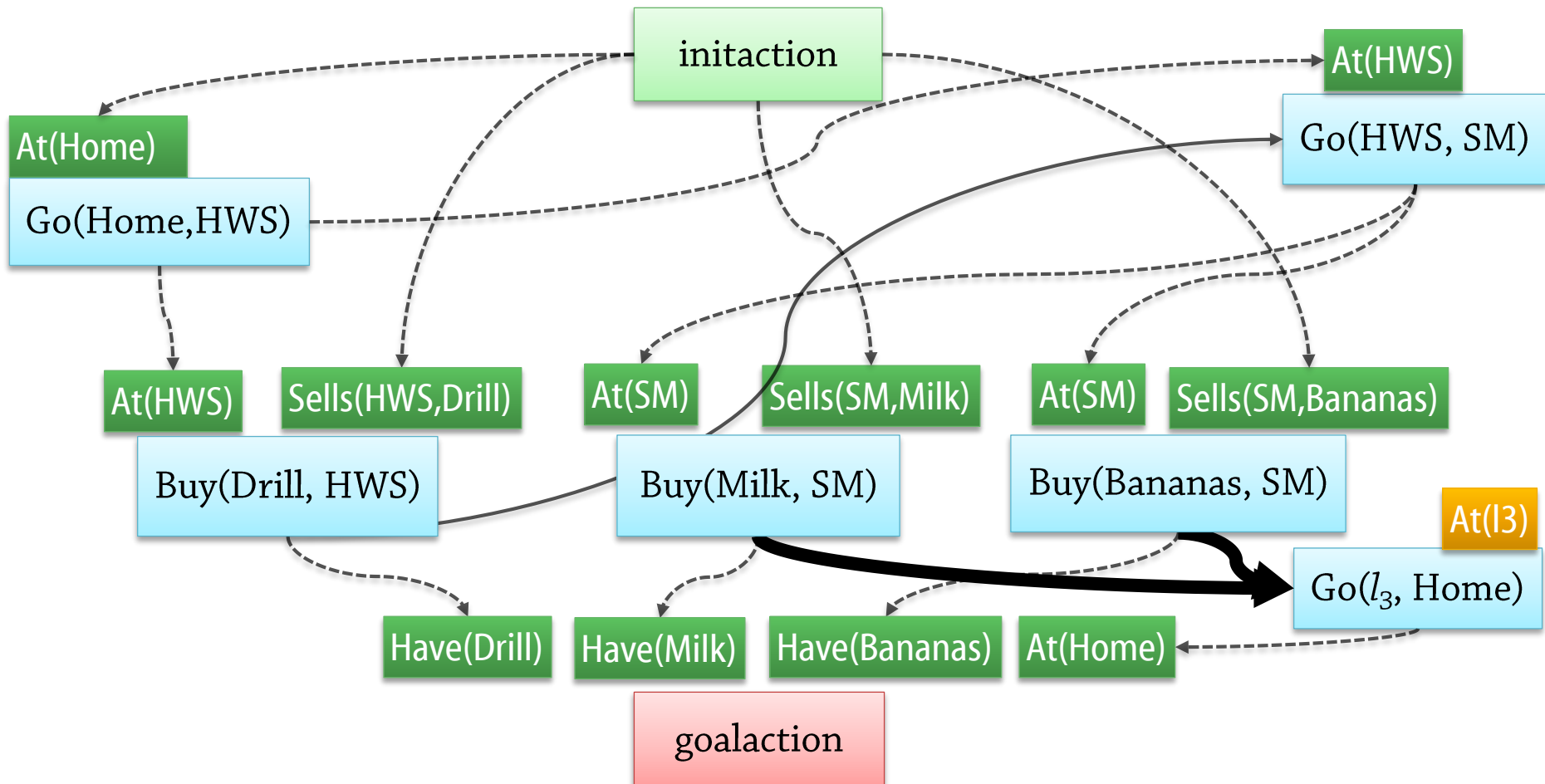
Example (continued)

- The only possible way to establish $\text{At}(\text{Home})$ for goalaction
 - This creates several new threats



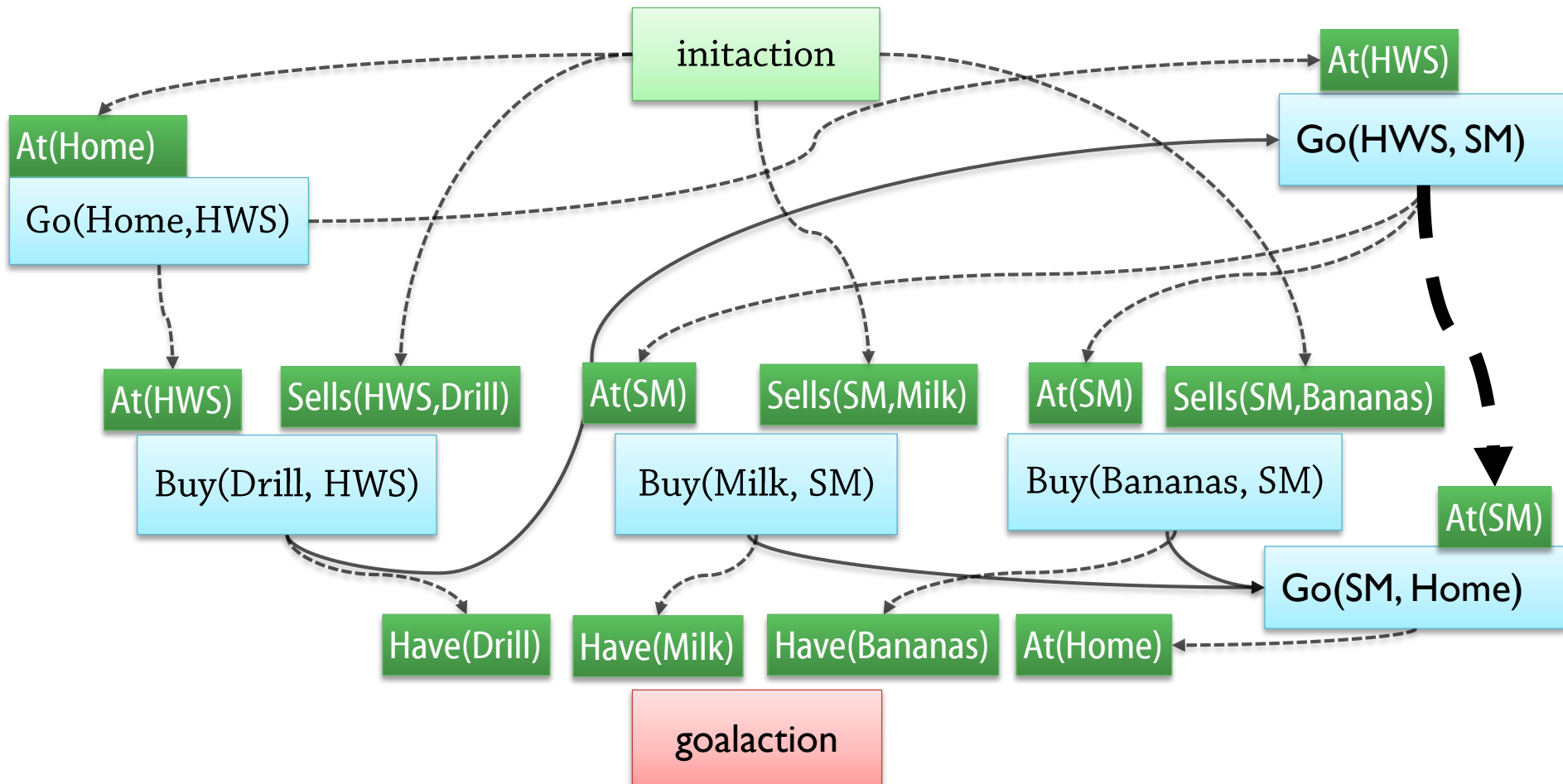
Example (continued)

- To remove the threats to $At(SM)$ and $At(HWS)$:
 - Make $go(HWS, SM)$ and $go(Home, HWS)$ precede $Go(l_3, Home)$
 - This also removes the other threats

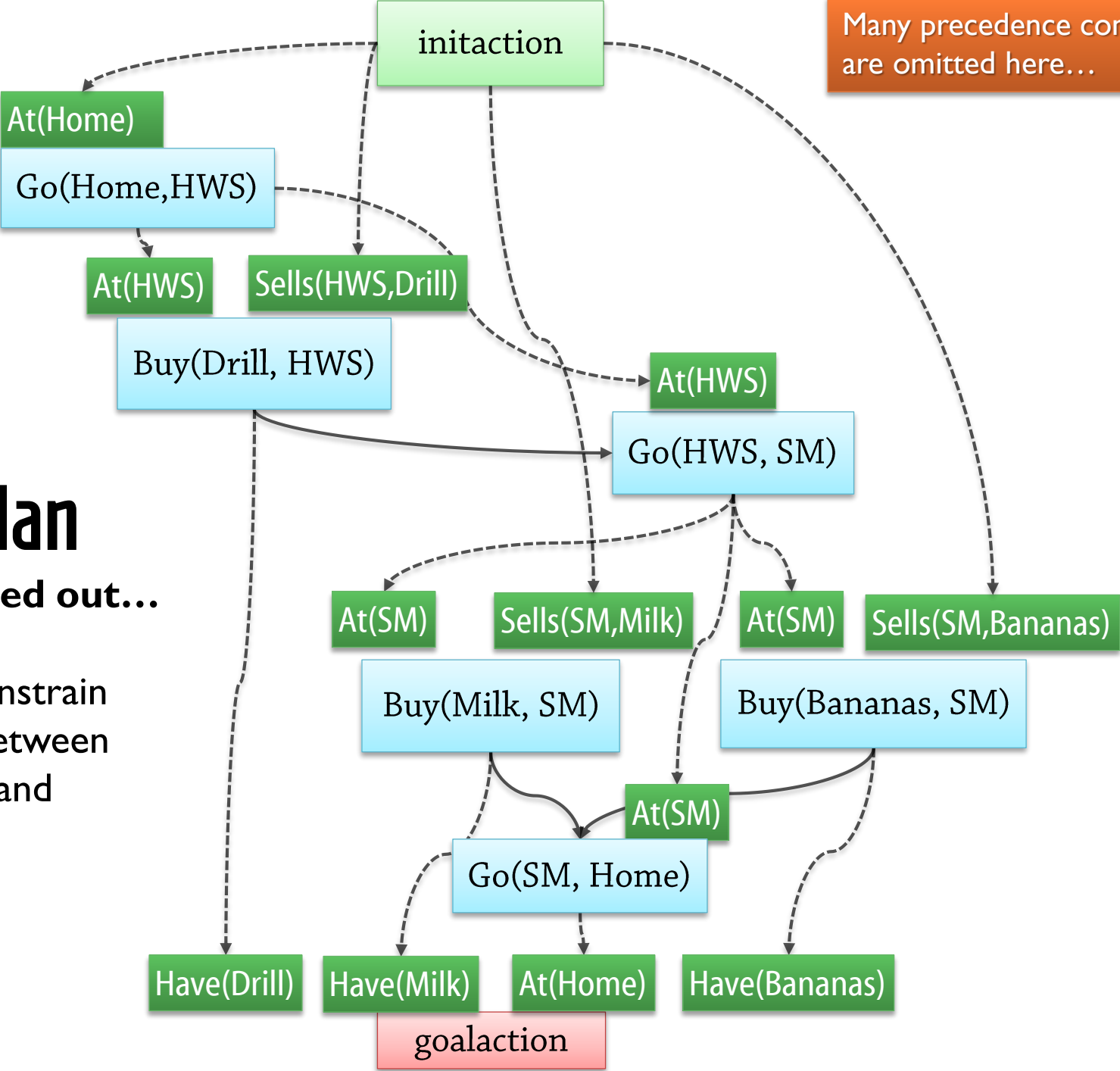


Final Plan

- Establish $At(l_3)$ with $l_3=SM$



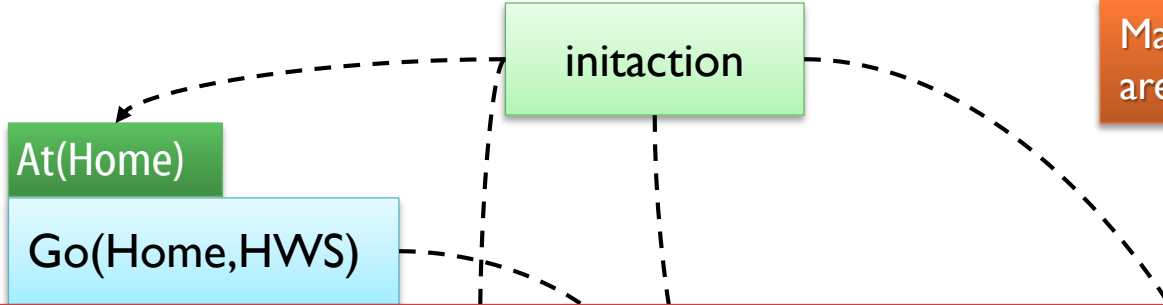
Many precedence constraints are omitted here...



Final Plan

Straightened out...

(Note: Still does not constrain the order between buying milk and bananas)

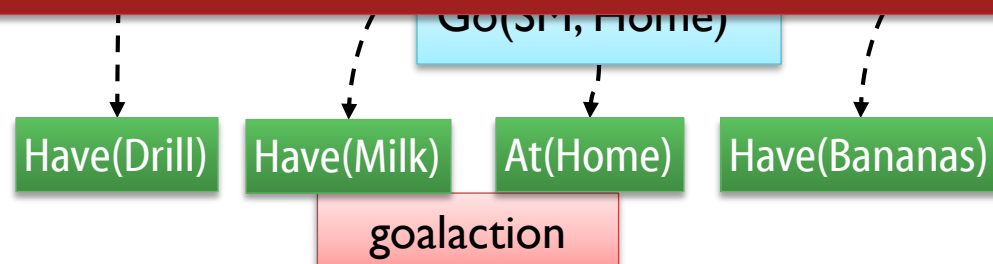


Many precedence constraints are omitted here...

This sequence assumed optimal choices!

Heuristics are required

Still, planners try many other alternatives, dead ends, etc.



Heuristics:

A few brief examples

■ POCL planning – one possible formulation (sound/complete):

```
■ planner(initial, goal):  
  open  $\leftarrow$  { Make-initial-plan(initial, goal) }  
  while open  $\neq \emptyset$ :  
    select a plan  $\pi$  in open  
    open  $\leftarrow$  open  $- \{ \pi \}$   
    flaws  $\leftarrow$  OpenGoals( $\pi$ )  $\cup$  Threats( $\pi$ )  
    if flaws =  $\emptyset$  then  
      return  $\pi$   
    else  
      select any flaw  $\varphi \in$  flaws  
      resolvers  $\leftarrow$  FindResolvers( $\varphi$ ,  $\pi$ ) // May be the empty set!  
      foreach r in resolvers:  
         $\pi' \leftarrow$  Refine( $\rho$ ,  $\pi$ ) // Actually apply the resolver  
        open  $\leftarrow$  open  $\cup \{ \pi' \}$   
  endwhile  
  return failure
```

- Examples of (generally non-admissible) **plan selection** heuristics:

- Count the **number of flaws** (see SNLP, UCPOP)

- $h_f(\pi) = |\{\text{flaws in } \pi\}|$

- Count the **number of open goals**

- $h_{oc}(\pi) = |\{\text{open goals in } \pi\}|$

- Adapt the **additive heuristic** (VHPOP)

$$h_{add}(\pi) = \sum_{\{open\ goals \xrightarrow{q} a_i \text{ in } \pi\}} h_{add}(q)$$

- Can be modified to address **reuse of actions** already in the plan (VHPOP):

$$h_{add}^r(\pi) = \sum_{open\ goals \xrightarrow{q} a_i \text{ in } \pi} \begin{cases} 0 & \text{if there is an action } a_j \in \pi \text{ that achieves } q \\ & \text{and could be before } a_i \\ h_{add}(q) & \text{otherwise} \end{cases}$$

- Adapt **planning graphs** (RePOP)

- Examples of flaw selection heuristics:
 - Choose **threats** before open goals (SNLP, UCPOP)
 - LIFO: Prefer the threat (or open goal) that was added last
 - Defer some threats until later
 - DSep: Delay threats that can be resolved through variable bindings ("separation")
 - DUnf: Delay threats that can be resolved in more than one way
 - Prefer the flaw with the fewest refinement options
 - Prefer open goals that must be resolved against the initial state
 - Cannot be achieved by any operator in the domain
 - As in plan selection, estimate costs of resolving flaws
 - $h_{\text{add}}()$, ...
 - ...

- Partial-order planning **delays commitment** to action ordering
 - Lower branching factor
 - More efficient in some situations
- Many POP planners still **assume sequential execution**
 - The intention was to find plans quickly,
not to find partially constrained plans