



ICAPS 2025 Tutorial

Learning for Generalised Planning

Dillon Z. Chen Felipe Trevizan Sylvie Thiébaux



<https://l4p-tutorial.github.io/>

Who is this tutorial for?

- planners
 - interested in “learning” for planning
- machine learners
 - interested in learning for “planning”
- knowledge representationers
 - interested in inductive reasoning in “learning for planning”

Assumed knowledge

- planning:
 - find a course of actions (i.e. a plan) or policy to achieve a goal
 - lifted representations, e.g. PDDL
- machine learning:
 - make predictions based on data and experience
 - training vs. testing pipelines

Agenda – Learning for Generalised Planning (L4P)

1. Introduction
2. Theory
 - 2.1. The L4P problem setup
 - 2.2. Methodologies using learning for L4P
 - 2.2.1. Forms of learned knowledge: policies, heuristics, transformations
 - 2.2.2. Tools for learning knowledge: graph learning, abstraction, constraint programming, policy search, language models
 - 2.3. Evaluating L4P methodologies
 - 2.3.1. Theoretical metrics: expressivity, generalisation, complexity
 - 2.3.2. Empirical metrics: training and planning costs, solution quality
3. Lab
 - 3.1. Getting used to the L4P setup: datasets and visualising data
 - 3.2. Testing and benchmarking existing L4P architectures
 - 3.3. Building a L4P architecture from scratch

Tutorial website
(contains links to related resources and these slides)

<https://l4p-tutorial.github.io/>

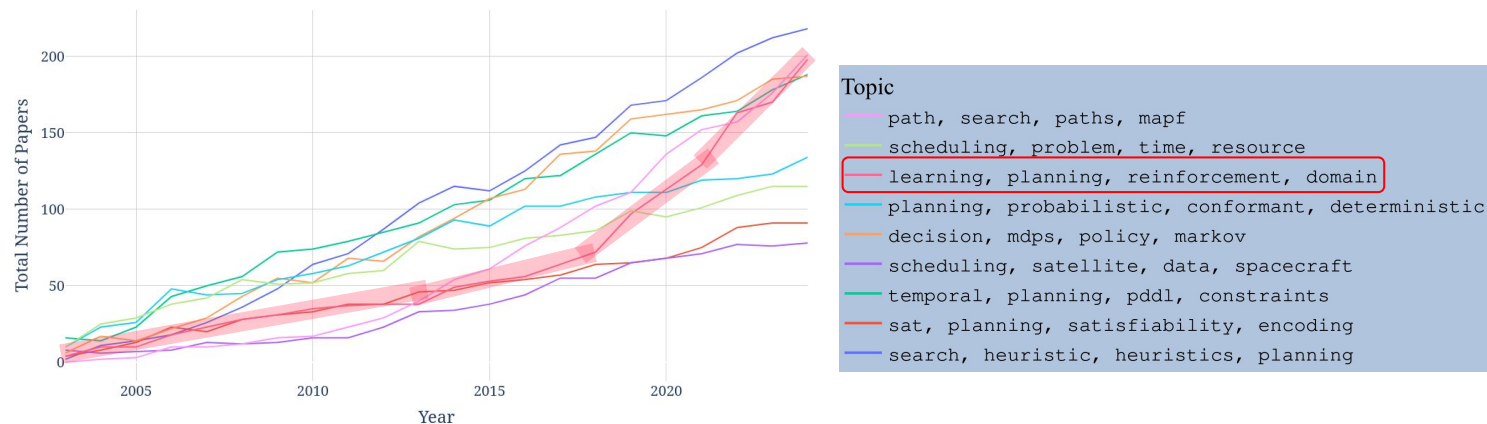
Article link
(comprehensive version of slides with more references)

Introduction

Motivation

Learning for (Generalised) Planning is a rapidly growing topic

- learn knowledge from easy-to-solve, small problems
- generalise to problems with unseen initial states/goals, and greater number of objects



Number of papers grouped by topic at ICAPS 2003–2024

Learning for Generalised Planning?

- Q: Can we leverage learning research for planning?
- Q: Does learning for planning work?
- Q: Is there still need for domain-independent planning research?

The diagram illustrates the AI planning process, showing the flow from a **planning problem** to a **planning model**, then to a **solver**, and finally to a **solution**.

planning problem: Represented by a visual of a block world with blocks of different colors (red, blue, green, yellow, orange) stacked on a table.

planning model: Shown as a code snippet defining the initial state and actions for a block world problem.

```
(define (domain blocksworld)
  (:domain blocksworld)
  (:requirements :strings)

  (:predicates (clear ?x ?w ?z ?u)
    (on-table ?x ?w ?z ?u)
    (arm-empty) ?w ?z ?u
    (holding ?x ?w ?z ?u)
    (on ?w ?y ?u ?z ?u))

  (:action pick-up
    (parameters ?w0)
    (precondition (and (clear ?w0) (on-table ?w0) (arm-empty)))
    (effect (and (holding ?w0) (not (clear ?w0)) (not (on-table ?w0))
      (not (arm-empty))))

  (:action put-down
    (parameters ?w0)
    (precondition (holding ?w0))
    (effect (and (clear ?w0) (arm-empty) (on-table ?w0)
      (not (holding ?w0))))

  (:action stack
    (parameters ?w0 ?w1 ?w2 ?w3)
    (precondition (and (clear ?w2 ?w3) (holding ?w0)
      (effect (and (arm-empty) (clear ?w0) (on ?w0 ?w2 ?w3)
        (not (clear ?w2 ?w3)) (not (holding ?w0))))))

  (:action unstack
    (parameters ?w0 ?w1 ?w2 ?w3)
    (precondition (and (on ?w0 ?w2 ?w3) (clear ?w0) (arm-empty)))
    (effect (and (holding ?w0) (clear ?w2) (not (clear ?w0)) (not
      (arm-empty))))))
```

solver: Represented by a code snippet showing the goal and the sequence of actions to solve the problem.

```
(define (problem blocksworld-01)
  (:domain blocksworld)
  no commands
  (:objects b1 b2 b3 b4 b5 - object)
  no commands
  (:init
    (arm-empty)
    (clear b3)
    (on b3 b5)
    (on b5 b4)
    (on-table b4)
    (clear b2)
    (on b2 b1)
    (on-table b1))
  (:goal (and
    (clear b4)
    (on b4 b3)
    (on-table b3)
    (clear b2)
    (on-table b2)
    (clear b1)
    (on b1 b5)
    (on-table b5))))
```

solution: Shown as a list of actions to be performed, such as `(unstack b5 b4)`, `(putdown b5)`, `(pickup b4)`, `(stack b4 b3)`, `(unstack b2 b1)`, `(putdown b2)`, `(pickup b1)`, and `(stack b1 b5)`.

Below the flow, two learning goals are indicated:

- Learn a model of the world**: An arrow points from this goal to the **planning model**.
- Learn to solve the model**: An arrow points from this goal to the **solver**.

What do we want with Learning for Generalised Planning?

*“learning generalised plans from example problems with the aim of
amortising the cost of learning by solving problems more efficiently (faster) and effectively
(higher quality solutions) than domain-independent planners”*

Preliminaries: Lifted Planning Problem

A lifted planning **domain** $\mathbf{D}$ is a tuple $\langle \mathbf{P}, \mathbf{A} \rangle$

- $\mathbf{P}$ = **predicates**; e.g. (at ?obj ?loc)
- $\mathbf{A}$ = **action schemata**; e.g. (move ?obj ?loc_from ?loc_to)

A lifted planning **problem** is a tuple $\langle \mathbf{D}, \mathbf{O}, \mathbf{s}_0, \mathbf{g} \rangle$

- $\mathbf{O}$ = **objects**; e.g. dog, kitchen
- $\mathbf{s}_0$ = **initial state**; e.g. (at dog bedroom), (hungry dog)
- $\mathbf{g}$ = **goal condition**; e.g. (happy dog)

ICAPS 2025 Tutorial

Learning for Generalised Planning: Theory

Dillon Z. Chen Felipe Trevizan Sylvie Thiébaux

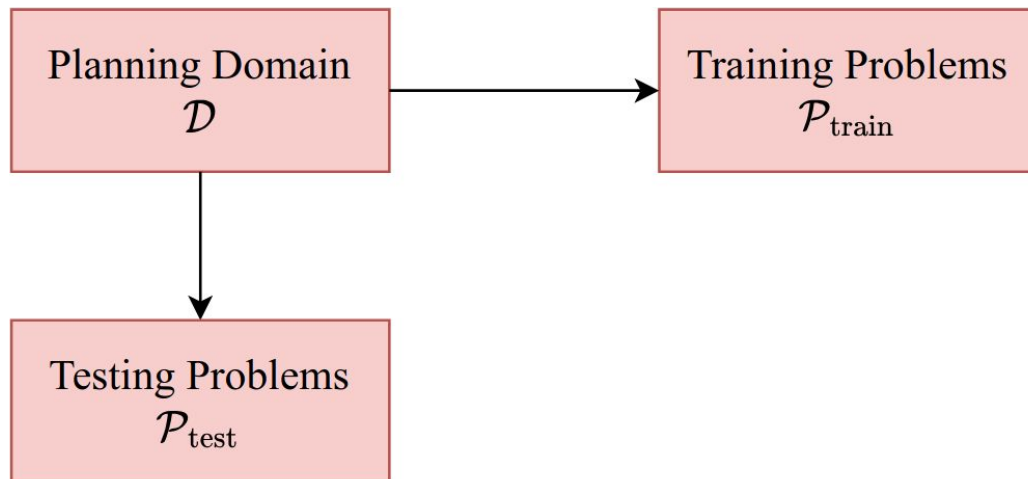


<https://l4p-tutorial.github.io/>

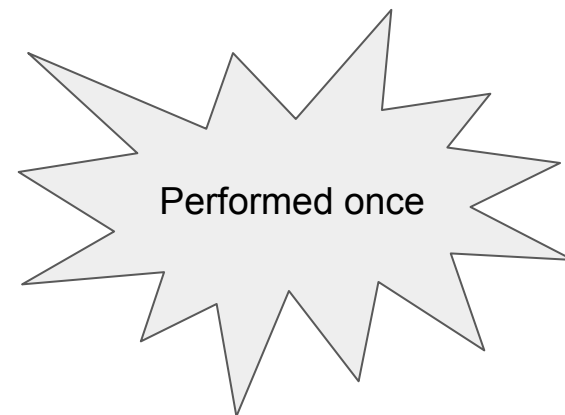
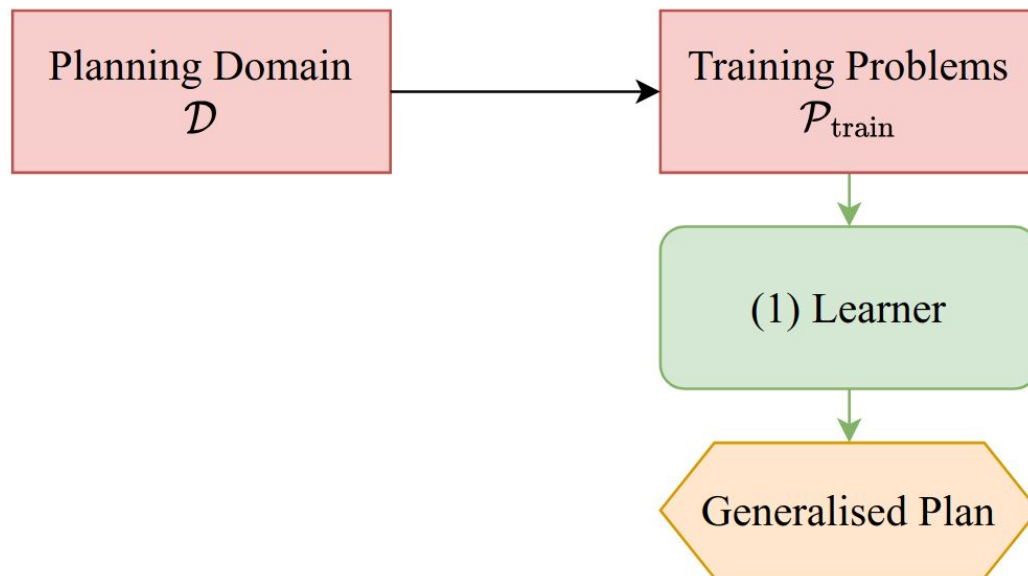
(2.1) The L4P problem setup

Learning for Generalised Planning (L4P) Setup

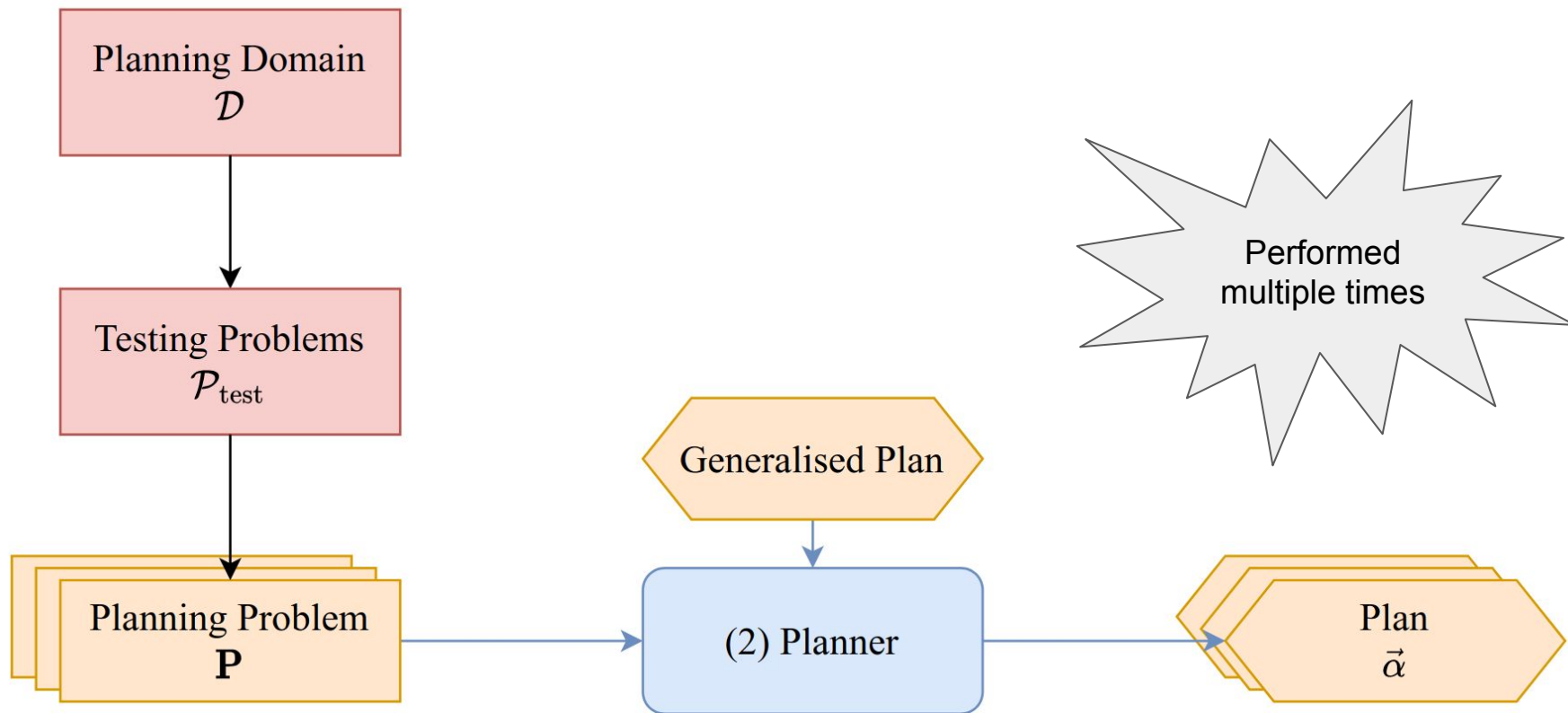
L4P Setup – Inputs



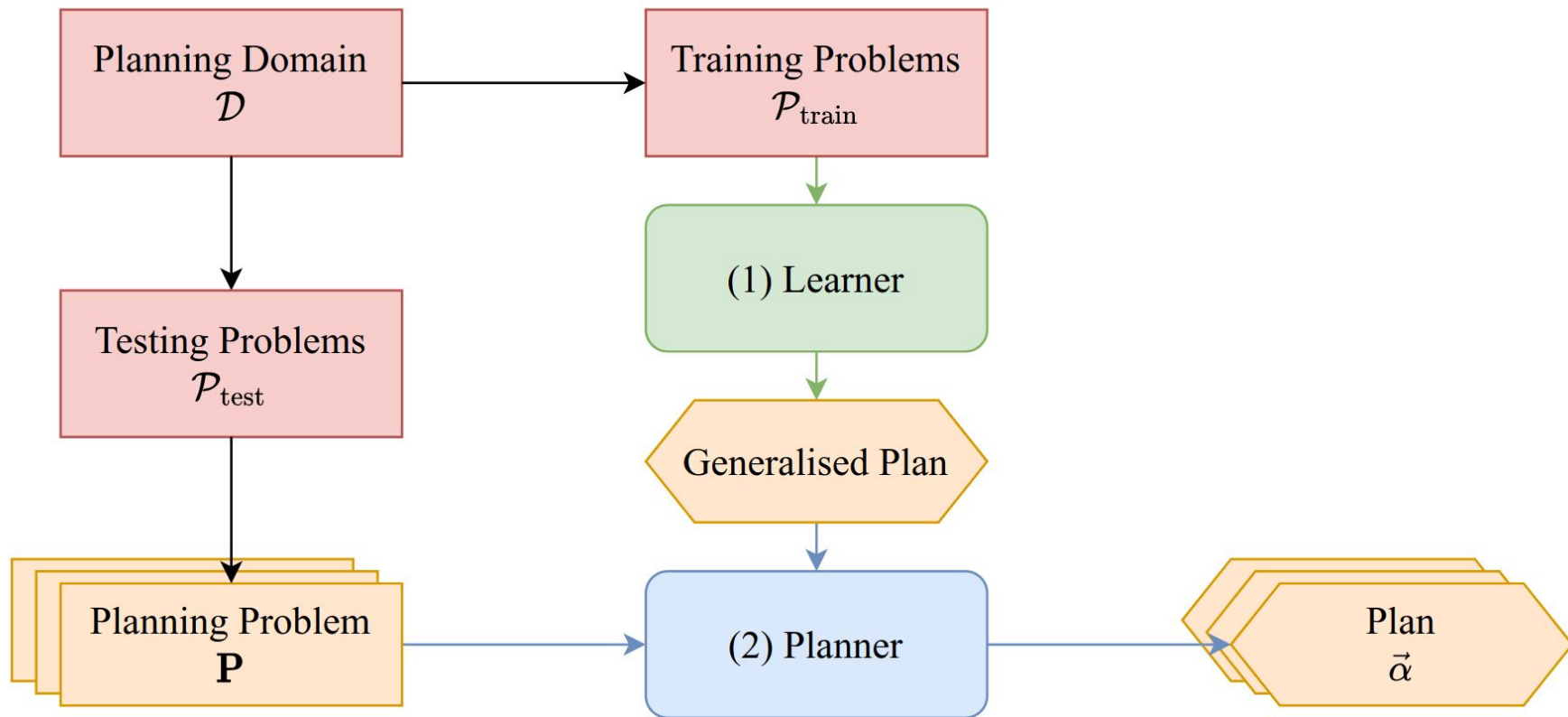
L4P Setup – Step 1: Learning



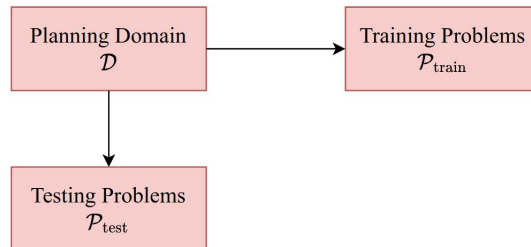
L4P Setup – Step 2: Planning



L4P Setup



Training vs Testing distributions



Define “ $f(x)$ = maximum number of objects and plan length of problems in x ”

1. Interpolation; in-distribution learning:

$$f(\mathcal{P}_{\text{test}}) \leq f(\mathcal{P}_{\text{train}})$$

2. Extrapolation; out-of-distribution learning:

$$f(\mathcal{P}_{\text{test}}) > f(\mathcal{P}_{\text{train}})$$

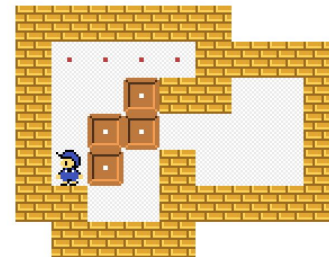
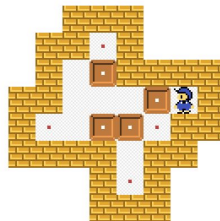
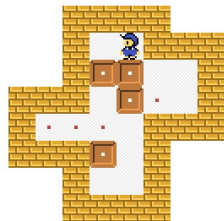
Training vs Testing distributions

Define “ $f(x)$ = maximum number of objects and plan length of problems in x ”

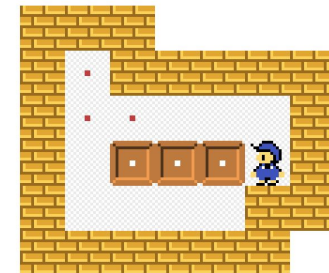
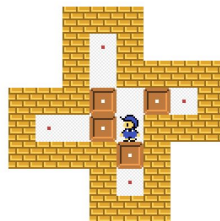
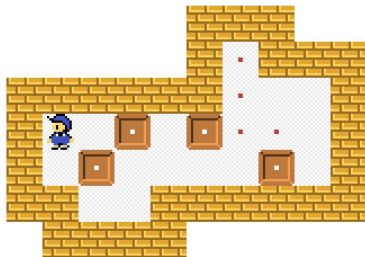
1. **Interpolation:** $f(\mathcal{P}_{\text{test}}) \leq f(\mathcal{P}_{\text{train}})$
 - for intractable domains
 - e.g. Sokoban, Rubik's cube; optimal planning
2. **Extrapolation:** $f(\mathcal{P}_{\text{test}}) > f(\mathcal{P}_{\text{train}})$
 - for tractable domains
 - e.g. package transportation, Blocksworld

GP interpolation example; Sokoban

Training

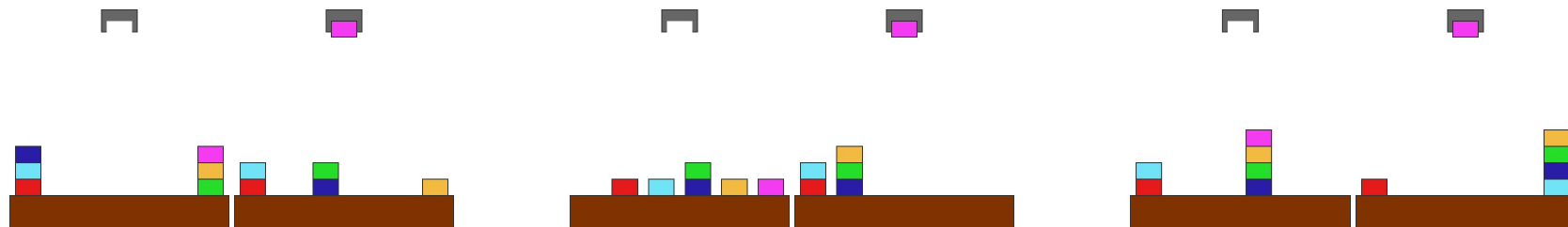


Testing: initial states and goals not seen before in training

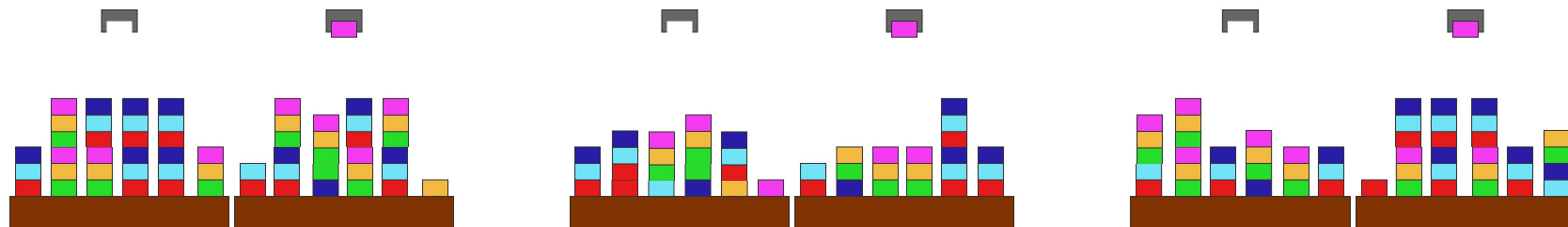


GP extrapolation example; Blocksworld

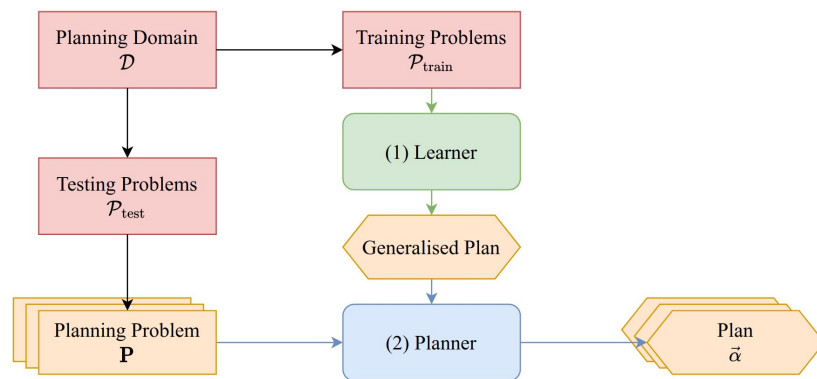
Training



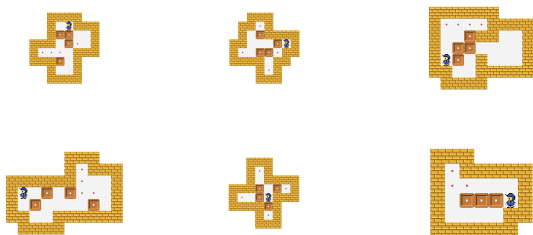
Testing: more blocks than seen in training problems



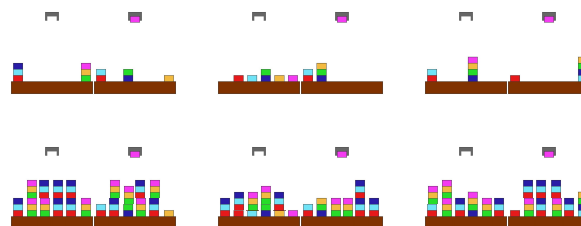
(2.1) The GP problem setup – *Summary*



1. Interpolation; in-distribution learning: $f(\mathcal{P}_{\text{test}}) \leq f(\mathcal{P}_{\text{train}})$

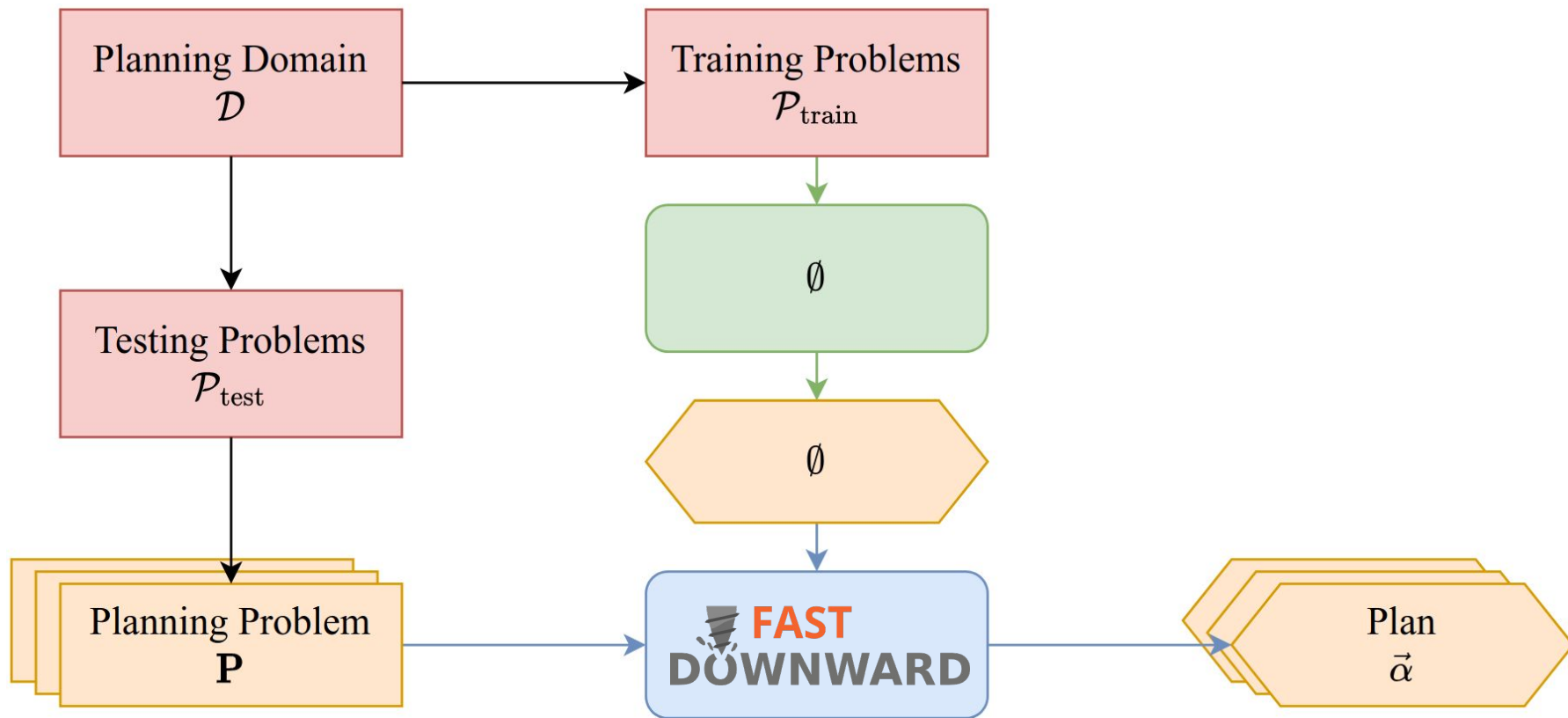


2. Extrapolation; out-of-distribution learning: $f(\mathcal{P}_{\text{test}}) > f(\mathcal{P}_{\text{train}})$



(2.2) Methodologies for solving GP problems

Trivial example: domain-independent planner with no learning



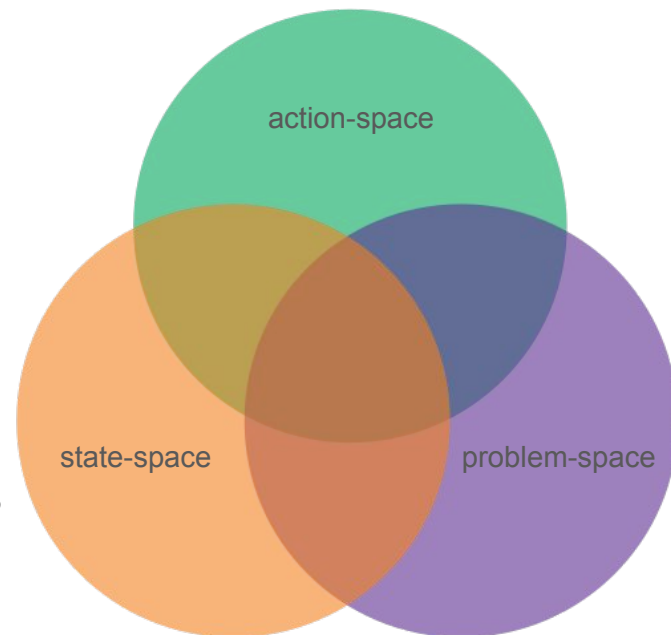
(2.2.1) Forms of learned knowledge

GP Knowledge Taxonomy

Learned knowledge can be taxonomised in one of 3 main spaces:

1. action-space; e.g. policies $\pi(a \mid s, g)$, finite state automata
2. state-space; e.g. heuristic $h(s)$, LTL constraints
3. problem-space; e.g. problem transformation

~ a method can be a combination of one or more spaces



Action-space methods: Policies

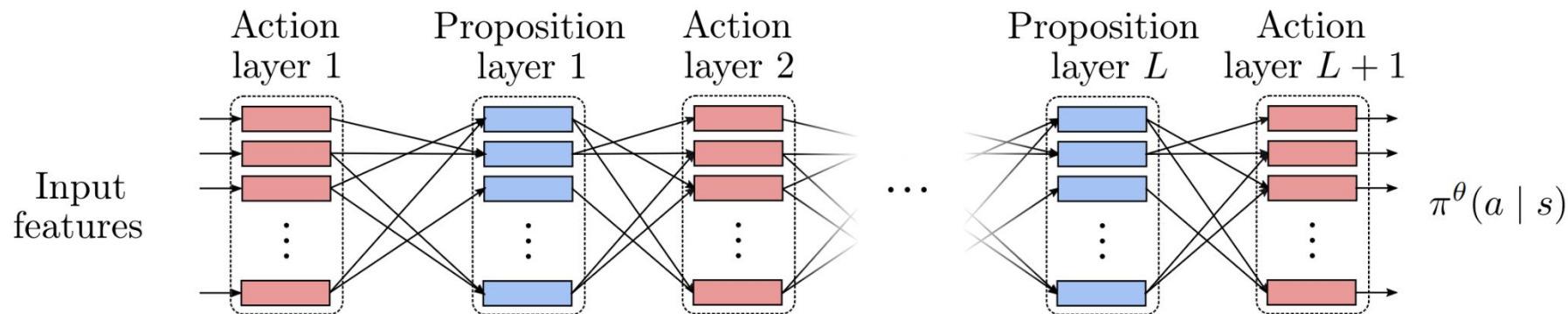
Action-space methods: Policies

- A policy, usually denoted $\pi(a \mid s)$, is a mapping from states to distributions of actions
- Learning
 - supervised: labelled optimal actions
 - reinforcement: improve based on signal
- Planning
 - usually rollout by sampling or choosing the argmax from $\pi(a \mid s)$
 - repeat until goal is reached or a time or rollout limit is reached
- Analogous to policy gradient methods in RL, e.g. DDPG, TRPO, PPO, etc.

Action-space example: ASNets

ASNets – architecture

- *Input*: planning state, goal condition, and action applicability
- *Output*: stochastic policy (distribution of actions)
- *Backbone*: Graph Neural Network



ASNets – learning

- *Reinforcement-learning style*; Repeat the following:
 - a) sample states by rolling out parameterised policy π^θ
 - b) sample states by executing a teacher planner
 - c) compute best actions for all sampled states via a teacher planner
 - d) update π^θ based on best actions; loss = binary cross entropy
- ~ similar to RL: explore (a) and exploit (b), with reward signals computed by a teacher planner (c) for improving the incumbent policy (d)

ASNets – planning

1. Set $s \leftarrow$ the initial state of the problem
2. Repeat the following:
 - a. return plan to s if goal is reached
 - b. sample an action a or select action with highest probability from $\pi^\theta(a \mid s)$
 - c. $s \leftarrow$ apply action a at s

ASNets – key attributes

- *Extrapolates* to testing problems larger than training problems
- *Sound* (returned solutions are correct)
- *Not complete* (if a solution exists, it is found)
- Plans on *grounded* representations
- Originally designed for *probabilistic planning*
- Extended to *numeric planning* [Wang and Thiébaux, ICAPS'24]

ASNets – notable achievements

- First deep learning architecture for learning generalised policies (2018)
- Blocksworld
 - a) 25 training problems with 8-10 blocks
 - b) solves all testing problems with 35-50 blocks
- Outperforms (probabilistic) planners on various domains
- Sparse regularisation loss for *interpretable* models

State-space methods: Heuristics

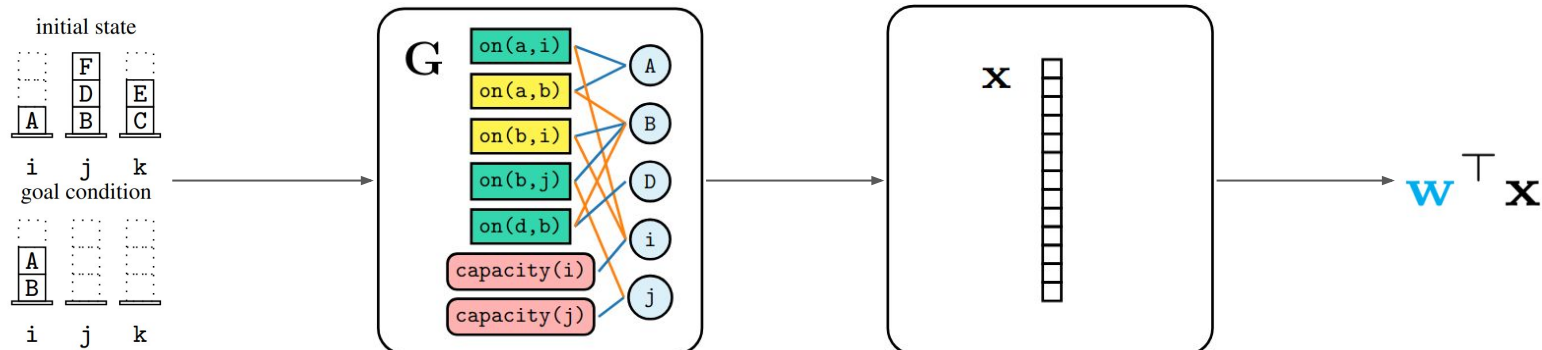
State-space methods: Heuristics

- A heuristic is a real-valued function $h: S \rightarrow \mathbf{R}$ estimating cost to go to a goal
- Learning
 - supervised: labelled optimal actions
 - reinforcement: improve based on signal
- Planning
 - combine with heuristic search, e.g. GBFS; but *NOT* A* as learned heuristics not guaranteed to be optimal
 - for greedy policies: select action whose successor has lowest h value
- Analogous to value function approximation in RL, e.g. Q-Learning, TD(λ), DQN, etc.

State-space example: (WL-)GOOSE

GOOSE – architecture

- *Input*: planning state and goal condition
- *Output*: scalar heuristic value
- *Backbone*: Weisfeiler-Leman Graph Kernel



GOOSE – learning

- *Supervised learning*
- Generate optimal plans from training data
- Optimise weights subject to predicting optimal
 - a) h^* values; loss = mean squared error
 - b) state rankings; loss = cf. [Garrett et al., IJCAI'16; Hao et al., IJCAI'24]

GOOSE – planning

- Plug into a heuristic search algorithm
 - e.g. GBFS

GOOSE – key attributes

- *Extrapolates* to testing problems larger than training problems
- *Sound* (returned solutions are correct)
- *Complete* (if a solution exists, it is found)
- Plans on *lifted* representations
- Extended to *probabilistic planning* [Zhang and Trevizan, Tech. Rep. 24]
- Extended to *numeric planning* [Chen and Thiébaux, NeurIPS'24]

GOOSE – notable achievements

- First learning approach to outperform domain-independent heuristics on competition IPC benchmarks with hundreds of objects (2024)
- Orders of magnitude more efficient than neural networks for learning and planning
- Maximally expressive compared to graph neural network counterparts
- Symbolic linear model with *interpretable* features

Problem-space methods: Problem transformation

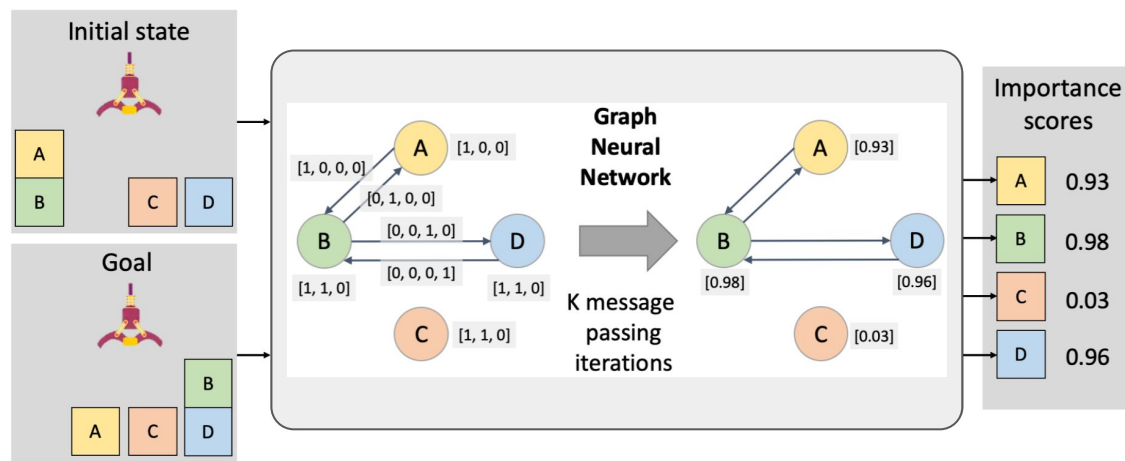
Problem-space methods: Problem transformation

- converts $Prob = \langle P, O, A, s_0, G \rangle$ into $Prob' = \langle P', O', A', s_0', G' \rangle$, ideally such that a plan for $Prob'$ can be transformed back to a plan for $Prob$
- advantage: scales with the improvement of planners over time

Problem-space example: PLOI

PLOI – architecture

- *Input*: planning state and goal condition
- *Output*: scalar scores in range $[0, 1]$ for each object in the problem
- *Backbone*: Graph Neural Network



PLOI – learning

- *Supervised learning*
- loss = binary cross entropy
- Greedily approximate ground truth labels of relevant objects in training data:
 - a) start with full object set O
 - b) incrementally remove some object from O until removing an object causes the problem to be unsolvable

PLOI – planning

1. score objects in a problem; select threshold value γ from the range $(0, 1)$
2. for $N = 1, 2, 3, \dots$ until the problem is solved:
 - a. keep the set of objects with score greater than γ^N
 - b. try to solve the problem with the subset of objects

PLOI – key attributes

- *Extrapolates* to testing problems larger than training problems
- *Sound* (returned solutions are correct)
- *Complete* (if a solution exists, it is found)
- Learns on *lifted* representations
- Can be extended to *probabilistic and numeric planning*

PLOI – notable achievements

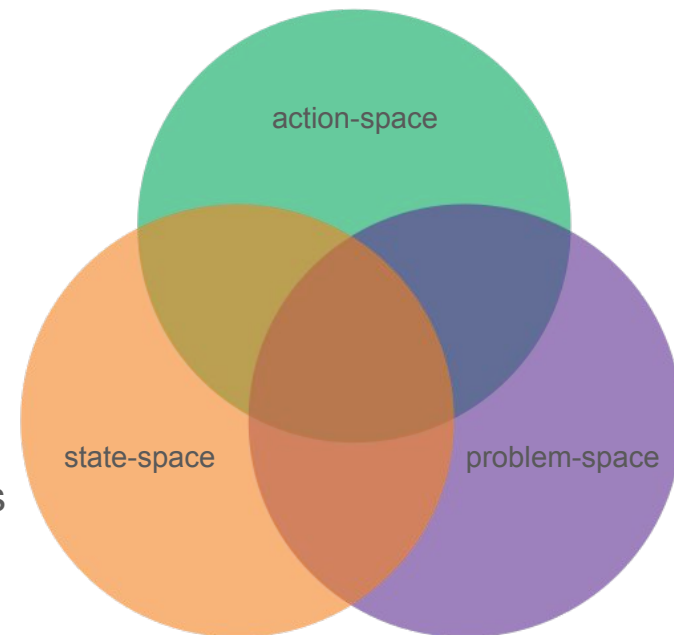
- Always matches or improves upon base planner in experiments
- Extended to handle continuous task and motion planning problems
- Learns a problem transformation $\Rightarrow$ can be used with any backend planner

(2.2.1) Forms of Learned Knowledge – *Summary*

Learned knowledge can be taxonomised in one of 3 main spaces:

1. action-space; e.g. policies
2. state-space; e.g. heuristic
3. problem-space; e.g. problem transformation

~ a method can be a combination of one or more spaces

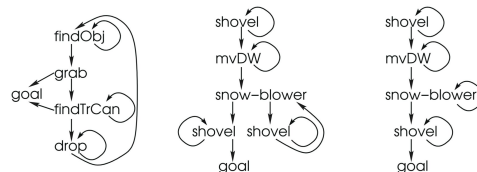


(2.2.2) Tools for learning knowledge

Approaches for synthesising generalised plans

Diverse body of approaches for synthesising generalised plans

- graph learning
- abstraction
- constraint programming
- program search
- goal regression
- language models

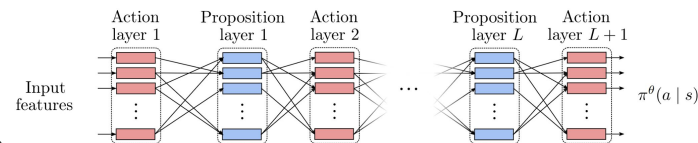


Srivastava et al., AAAI'11

Algorithm 1: Best-First Generalized Planning (BFGP).

Data: The GP problem P , pointers Z , program lines n , a list of evaluation functions P'
Result: A generalized plan Π that solves P
 $Open_P \leftarrow \{\Pi_{empty}\};$
while $Open_P \neq \emptyset$ **do**
 $\Pi \leftarrow \text{extractBestProgram}(Open_P);$
 $ChildrenPrograms \leftarrow \text{expandProgram}(\Pi, Z, n);$
 for $IV \in ChildrenPrograms$ **do**
 if $isDeadEnd(IV, P)$ **then**
 continue
 if $isGoal(IV, P)$ **then**
 return (Π')
 $Open_P \leftarrow \text{insertProgram}(Open_P, \Pi');$
 end
end
return $\emptyset$ // no solution found

Segovia-Aguas et al., ICAPS'21; AIJ'24



Toyer et al., AAAI'18; JAIR'20

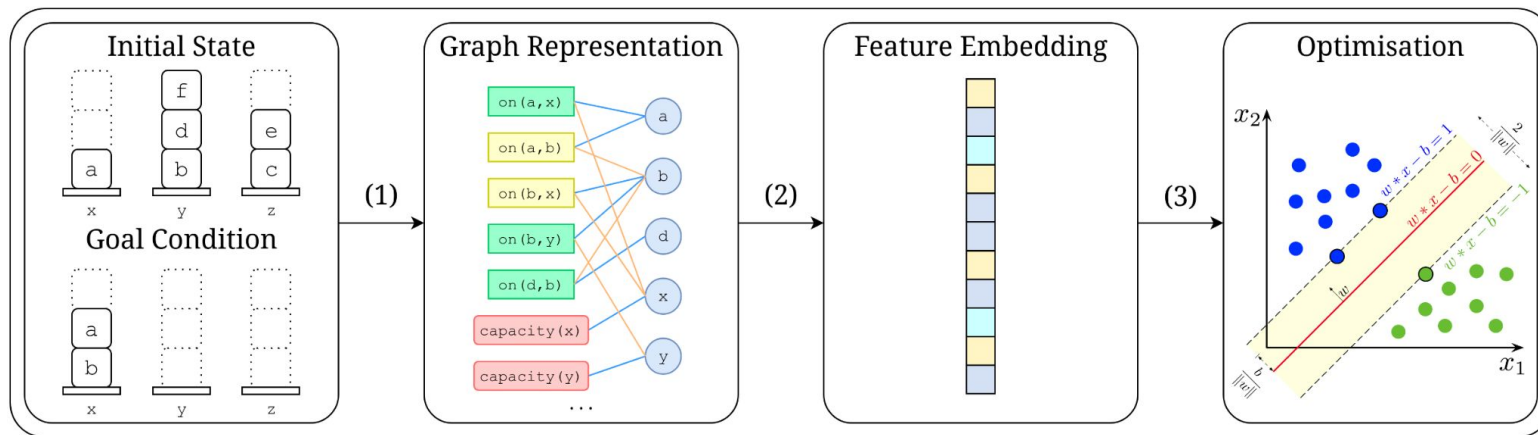
1. Policy: $\bigvee_{(s,s')} Good(s, s')$, s is non-goal state,
2. V_1 : Exactly-1 $\{V(s, d) : V^*(s) \leq d \leq \delta V^*(s)\}^1$
3. V_2 : $Good(s, s') \rightarrow V(s, d) \wedge V(s', d'), d' < d$,
4. Goal: $\bigvee_{f: [f(s)] \neq [f(s')]} Select(f)$, one $\{s, s'\}$ is goal,
5. Bad trans: $\neg Good(s, s')$ for s solvable, and s' dead-end,
6. D2-sep: $Good(s, s') \wedge \neg Good(t, t') \rightarrow D2(s, s'; t, t')$, where $D2(s, s'; t, t')$ is $\bigvee_{\Delta_f(s, s') \neq \Delta_f(t, t')} Select(f)$.

Francès et al., AAAI'21



Graph Learning

e.g. ASNets, PLOI, GOOSE in Sec. 2.2.1.



representation

- domain information: predicates, schema
- problem information: objects, state, goals

architecture

- graph neural networks
- graph kernels
- transformers (GNNs + positional encodings)

optimisation

- imitation learning, supervised, unsupervised, reinforcement learning
- MSE, CE, ranking

Abstraction

- e.g. place all the dirty objects on the floor into the laundry basket

Day 1

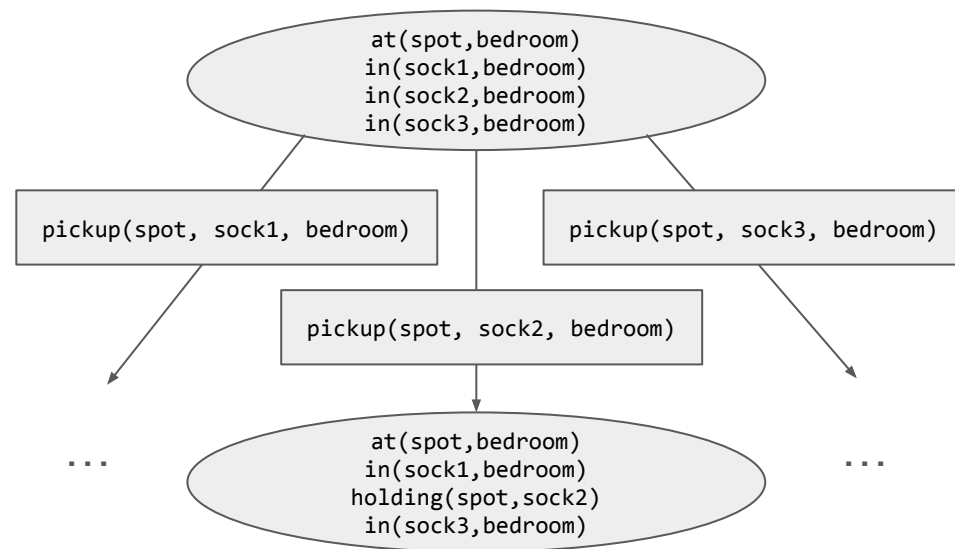
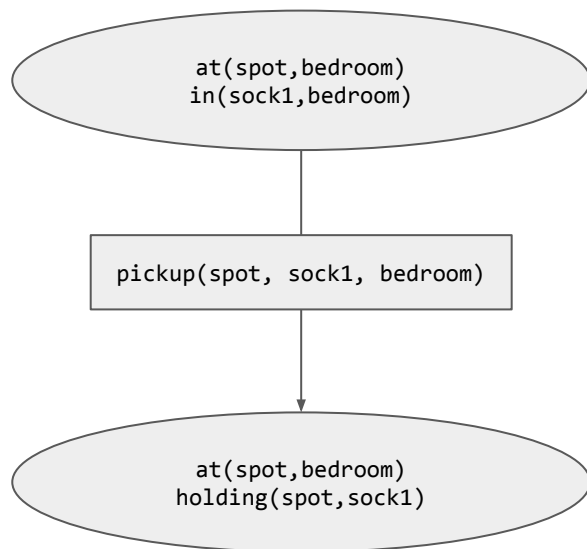


Day 10



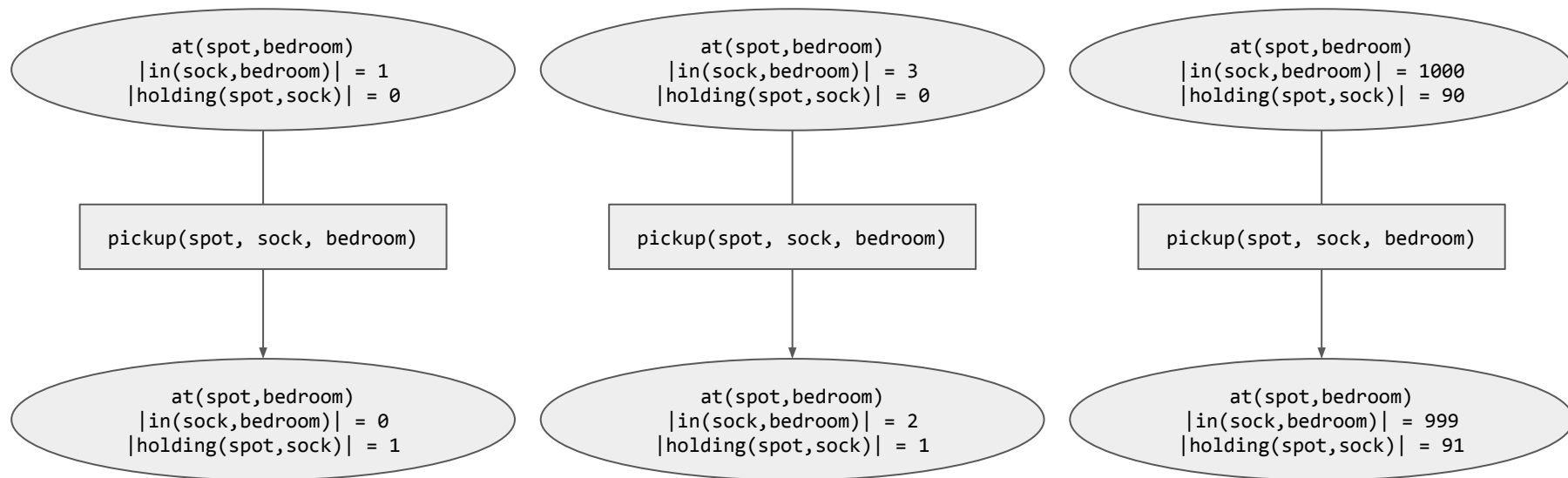
How to synthesise a policy across instances?

- many different instances = many different transitions, ... what can we do?



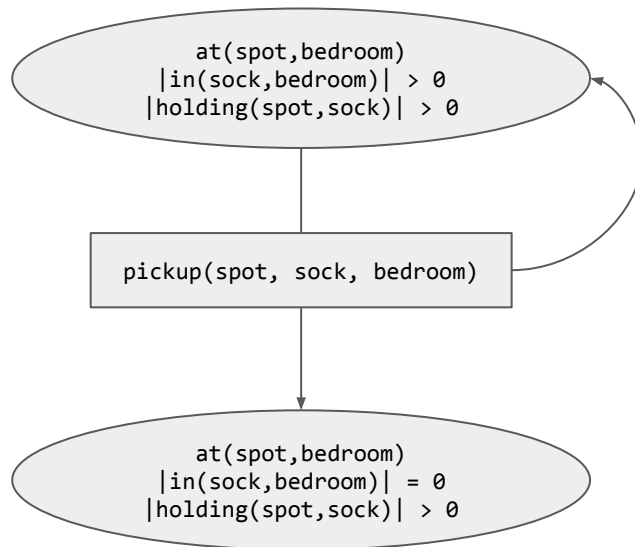
Abstraction

- we can abstract equivalent objects away:

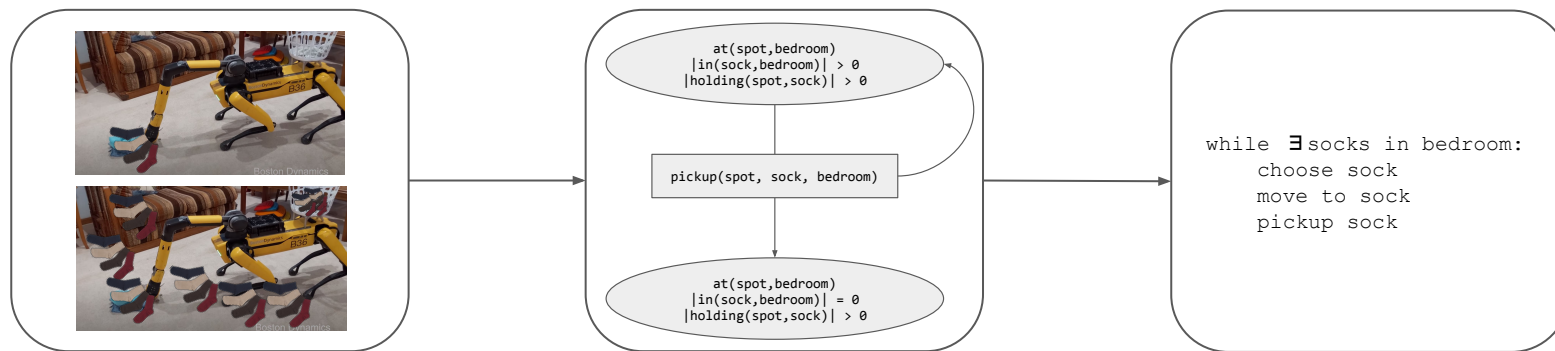


Abstraction

- we can even abstract numbers away via **qualitative states** [Srivastava et al., AAAI'11]



Abstraction



find an abstraction

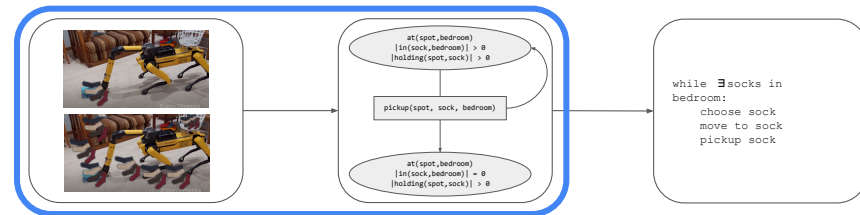
- synthesise a single nondeterministic problem representing a family of problems

solve the abstraction

- synthesise a generalised plan for the original family of problems from solving the abstraction

Finding Abstractions

- Illanes and McIlraith, AAI'19
 - bag equivalent objects



[Riddle et al., HSDIP'16, Fuentetaja and Rosa, AI Comm. 2016]

- Bonet et al., AAI'19
 - MaxSAT over description logic features for planning

[Martín and Geffner, Appl. Intell. 2004]

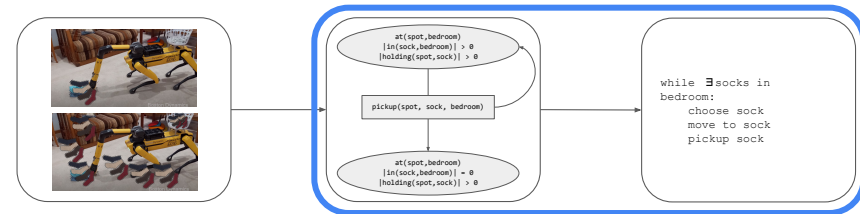
Solving Abstractions

- ‘generate-and-test’

- iteratively *generate* policies and *test* for termination
- generate: e.g. anytime FOND solver [Illanes and McIlraith, AAI'19; Zeng et al., IJCAI'22]
- test: e.g. Sieve algorithm [Srivastava et al., AAI'11; Srivastava, JAIR'23]

- compilation approach

- compile to FOND⁽⁺⁾ restricted to terminating policies and solve
- FOND: [Bonet and Geffner, JAIR'20]; FOND⁺: [Rodriguez et al., JAIR'22]



Constraint Programming

- finding and solving abstractions ‘end-to-end’
- synthesise generalised policies with SAT [Francès et al., AAAI'21]
- synthesise generalised policy sketches with ASP [Drexler et al, ICAPS'22]

- Policy: $V_{(s,s')} \text{ Good}(s, s')$, s is non-goal state,
- V_1 .: Exactly-1 $\{V(s, d) : V^*(s) \leq d \leq \delta V^*(s)\}$,¹
- V_2 : $\text{Good}(s, s') \rightarrow V(s, d) \wedge V(s', d')$, $d' < d$,
- Goal: $\bigvee_{f: \llbracket f(s) \rrbracket \neq \llbracket f(s') \rrbracket} \text{Select}(f)$, one $\{s, s'\}$ is goal,
- Bad trans: $\neg \text{Good}(s, s')$ for s solvable, and s' dead-end,
- D2-sep: $\text{Good}(s, s') \wedge \neg \text{Good}(t, t') \rightarrow D2(s, s'; t, t')$,
where $D2(s, s'; t, t')$ is $\bigvee_{\Delta_f(s, s') \neq \Delta_f(t, t')} \text{Select}(f)$.

Listing 1: Full ASP code for learning sketches: constraints C1-C8 satisfied. Optimization in lines 24-25 for finding a simplest solutions measured by number of sketch rules plus sum of feature complexities.

```

1 # C1: construct rules and select features.
2 { select(F) :- feature(F).
3   e_eq(R, F); e_gt(R, F); e_unk(R, F) :- 1 :- rule(R), numerical(F).
4   e_pos(R, F); e_neg(R, F); e_unk(R, F) :- 1 :- rule(R), boolean(F).
5   e_dec(R, F); e_inc(R, F); e_unk(R, F); e_bot(R, F) :- 1 :- rule(R), numerical(F).
6   e_pos(R, F); e_neg(R, F); e_unk(R, F); e_bot(R, F) :- 1 :- rule(R), boolean(F).
7 # C4 and C7: good and bad state pairs must comply with rules and selected features.
8 { good(R, I, S, S') :- rule(R), a_distance(I, S, S', _).
9   e_satisfied(R, F, I, S, S') :- ( e_dec(R, F) :- V > 0; e_inc(R, F) :- V < V';
   e_pos(R, F) :- V' = 1; e_neg(R, F) :- V' = 0; e_bot(R, F) :- V = V'; e_unk(R, F) :- 1,
   rule(R), feature_valuation(F, I, S, V), feature_valuation(F, I, S', V'),
   a_distance(I, S, S', _).
11 :- ( not e_satisfied(R, F, I, S); not e_satisfied(R, F, I, S, S') ) :- 0, select(F), good(
   R, I, S, S').
12 :- ( not e_satisfied(R, F, I, S); select(F); not e_satisfied(R, F, I, S, S') ; select(F)
   ) :- 0, rule(R), a_distance(I, S, S', _), not good(R, I, S, S').
13 # C2: there must be one subgoal tuple for each state with unbounded width.
14 { subgoal(I, S, T) :- tuple(I, S, T) :- 1 :- solvable(I, S), exceed(I, S).
15 # C3: states underlying subgoal tuples must be good.
16 :- ( good(R, I, S, S') :- rule(R) :- 0, subgoal(I, S, T), contain(I, S, T, S').
17 # C5: good pairs to dead-end states must be at larger distance than subgoal tuple.
18 :- D <= D', a_distance(I, S, S', D), t_distance(I, S, T, D'), subgoal(I, S, T),
   good(I, S, S'), solvable(I, S), unsolvable(I, S').
19 # C6: ensure acyclicity.
20 order(I, S, S') :- solvable(I, S), solvable(I, S'), good(I, S, S'), order(I, S').
21 order(I, S) :- solvable(I, S), order(I, S, S') :- good(I, S, S'), solvable(I, S),
   solvable(I, S').
22 :- solvable(I, S), not order(I, S).
23 # Optimization objective: smallest number of rules plus the sum of feature complexities.
24 #minimize { C, complexity(F, C) :- complexity(F, C), select(F) }.
25 #minimize { 1, rule(R) :- rule(R) }.
```

Program Search

- directly search over space of programs/generalised plans
- compile program search as planning problem: [Segovia-Aguas et al., AIJ'19]
- Best-First Generalised Planning: GP as heuristic search [Segovia-Aguas et al., JAIR'24]
- programs consist of
 - planning actions
 - goto instructions
 - termination instructions

```
CORRIDOR
0. move_right(l1,l2)
1. set(l1,l2)
2. inc(l2)
3. goto(0,~(x2^~y2))
4. move_left(l1,l2)
5. set(l1,l2)
6. dec(l2)
7. test(goal_at(l1))
8. goto(4,~(y2^x2))
9. end
```

```
FIBONACCI
0. vector_add(i,j)
1. dec(i)
2. vector_add(i,j)
3. set(j,i)
4. inc(i)
5. goto(0,~(x2^~y2))
6. end
```

```
FIND
0. accumulate(i)
1. inc(i)
2. goto(0,~(x2^~y2))
3. end
```

```
GRIPPER
0. pick(b1,r1,g1)
1. inc(r2)
2. move(r1,r2)
3. drop(b1,r2,g1)
4. move(r2,r1)
5. inc(b1)
6. goto(0,~(x2^~y2))
7. end
```

```
REVERSE
0. set(i,j)
1. swap(i,j)
2. inc(i)
3. goto(1,~(x2^~y2))
4. inc(i)
5. goto(0,~(x2^~y2))
6. end
```

```
SELECT
0. cmp(vector(i),vector(j))
1. goto(3,~(x2^~y2))
2. set(j,i)
3. inc(i)
4. goto(0,~(x2^~y2))
5. select(j)
6. end
```

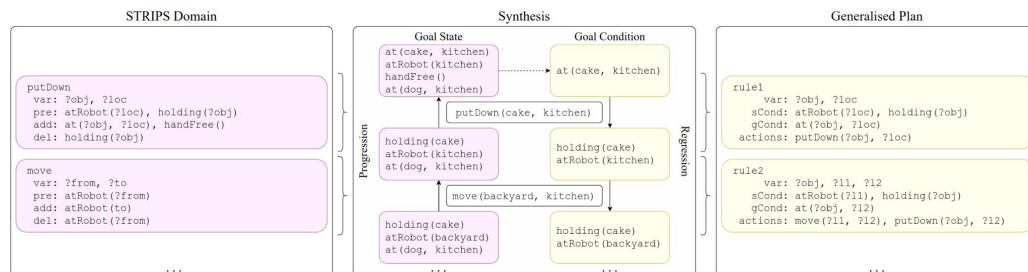
```
SORTING
0. cmp(vector(i),vector(j))
1. goto(5,~(x2^~y2))
2. swap(i,j)
3. dec(i)
4. dec(i)
5. set(j,i)
6. inc(i)
7. goto(0,~(x2^~y2))
8. end
```

```
TRIANGULAR SUM
0. vector_add(i,j)
1. set(j,i)
2. inc(c2)
3. goto(0,~(x2^~y2))
4. end
```

```
VISITALL
0. inc(c1)
1. move_right(c2,c1,r1)
2. inc(c2)
3. goto(0,~(x2^~y2))
4. inc(r1)
5. move_up(r2,r1,c1)
6. dec(c1)
7. move_left(c2,c1,r1)
8. dec(c2)
9. goto(5,~(x2^~y2))
10. inc(r2)
11. goto(0,~(x2^~y2))
12. end
```

Goal Regression

- goal regression determines **minimal and sufficient condition** for achieving a goal under a sequence of actions
- used to synthesise generalised, first-order policies by
 - Gretton and Thiébaux, UAI'04
 - Illanes and McIlraith, AAI'19
 - Chen et al., AAI'26



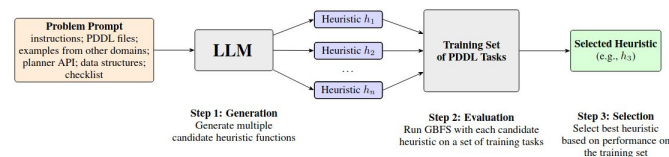
Language Models

prompt for programs representing generalised plans

```
def get_plan(objects, init, goal):
    # Helper function to extract objects of a specific type
    def get_objects_of_type(type_name):
        return [obj for obj in objects if obj.startswith(type_name)]
    # Helper function to find the initial or goal location of a ball
    def get_location(state, ball):
        for atom in state:
            if atom[0] == 'at' and atom[1] == ball:
                return atom[2]
        return None
    # Extract balls, rooms, and grippers from objects
    balls = get_objects_of_type('ball')
    rooms = get_objects_of_type('room')
    grippers = get_objects_of_type('gripper')
    ...
```

standalone solvers

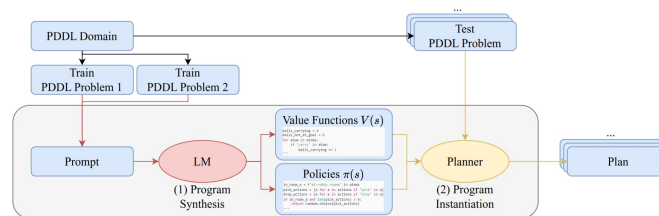
Silver et al., AAI'24



generalised heuristic functions

Tuisov et al., 2025

figure from Corrêa et al., NeurIPS'25



generalised action policies

Chen et al., EWRL'25

(2.3) Evaluating L4P Methodologies

Theoretical and Empirical Measures

- Theory
 - expressivity
 - generalisability
 - complexity and decidability
- Practice
 - training costs
 - planning costs
 - solution quality

Theoretical Measures – Expressivity

- expressivity = range of functions or hypotheses a model can potentially learn
- ~ measures the *best theoretical performance* of a model
- e.g.
 - domain-dependent GNNs related to C_2 logic [Stålberg et al., ICAPS'22]
 - domain-independent GNNs can express h^{max}, h^{add} [Chen et al., AAAI'24]

Theoretical Measures – Generalisability

- generalisability = performance on unseen data
- ~ measures the **estimated average performance** of a model
- see e.g. <https://mlstory.org/generalization.html>
 - Algorithmic stability
 - Vapnik–Chervonenkis (VC) dimension
 - Rademacher complexity
 - PAC learning
- almost no theory for generalised planning

Theoretical Measures – Complexity and Decidability

- theoretical computational resources required to synthesise a generalised plan
- usually trade-off with expressivity
- e.g.
 - QNP (planning with loops with non-deterministic semantics) is decidable (EXPTIME-complete) [Srivastava et al., AAAI'11; Bonet and Geffner, JAIR'20]
 - planning with loops with deterministic semantics is undecidable [Srivastava et al., AAAI'15; Srivastava, JAIR'23]

Empirical Measures

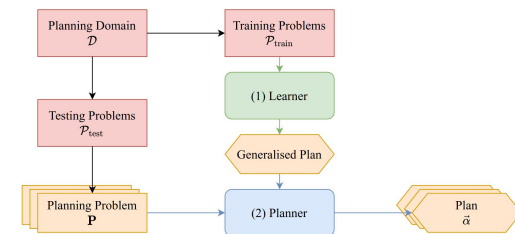
- Training Costs
 - resources for synthesising a generalised plan
- Planning Costs
 - resources for instantiating a generalised plan
- Solution Quality
 - quality (e.g. size) of a generalised plan, and of instantiated plans

Benchmarks

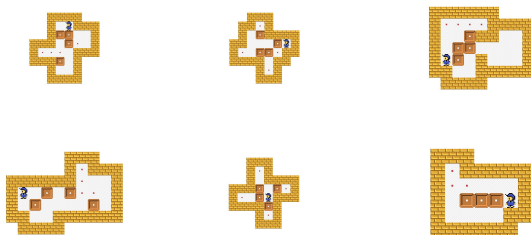
- Various papers often introduce their own train and test splits
- IPC 2023 Learning Track introduced canonical splits

<https://github.com/ipc2023-learning/benchmarks>

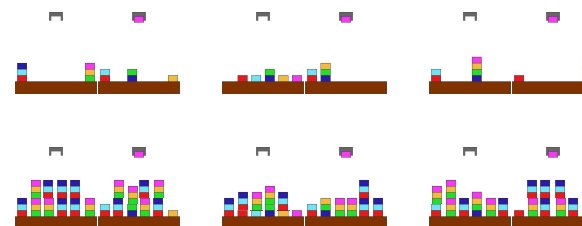
- Recall interpolation and extrapolation settings. Most L4P works focus on extrapolation



1. Interpolation; in-distribution learning: $f(\mathcal{P}_{\text{test}}) \leq f(\mathcal{P}_{\text{train}})$



2. Extrapolation; out-of-distribution learning: $f(\mathcal{P}_{\text{test}}) > f(\mathcal{P}_{\text{train}})$



ICAPS 2025 Tutorial

Learning for Generalised Planning: Lab

Dillon Z. Chen Felipe Trevizan Sylvie Thiébaux



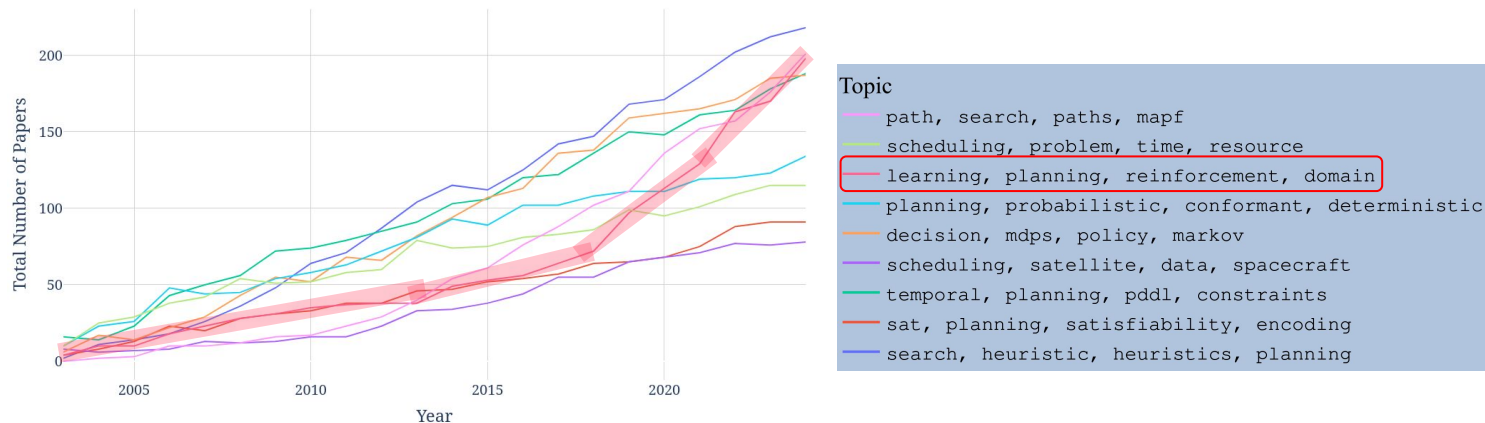
https://github.com/DillonZChen/wlplan/blob/main/docs/source/tutorials/1_introduction.ipynb

<https://l4p-tutorial.github.io/>

Motivation

Learning for (Generalised) Planning is a rapidly growing topic

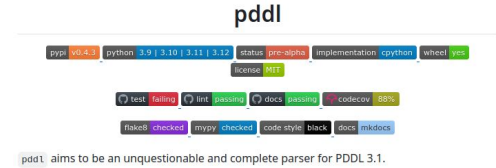
- learn knowledge from easy-to-solve, small problems
- generalise to problems with unseen initial states/goals, and greater number of objects



Number of papers grouped by topic at ICAPS 2003–2024

Motivation

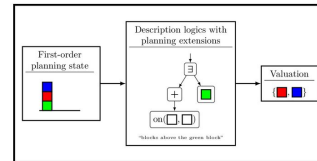
- plenty of open source planner systems and planning libraries



- few libraries specifically for learning for planning

Description Logics State Features for Planning Library (DLPlan)

Welcome to the DLPlan library.



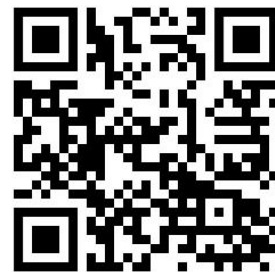
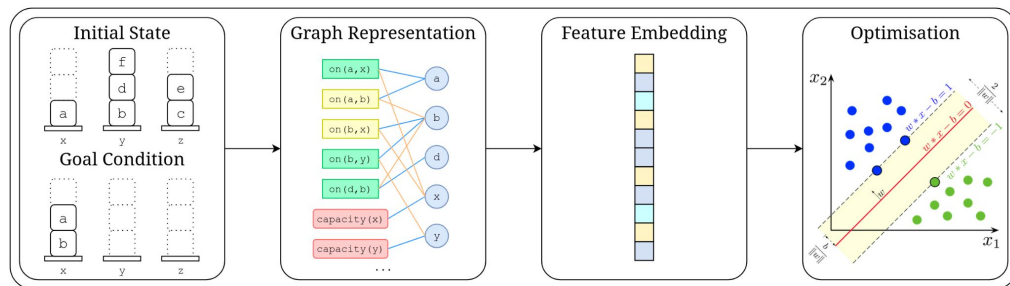
WLPlan

C++ package with Python bindings that implements

1. Graph representations of planning tasks
2. Embeddings of planning tasks and graphs
3. Serialisation of models

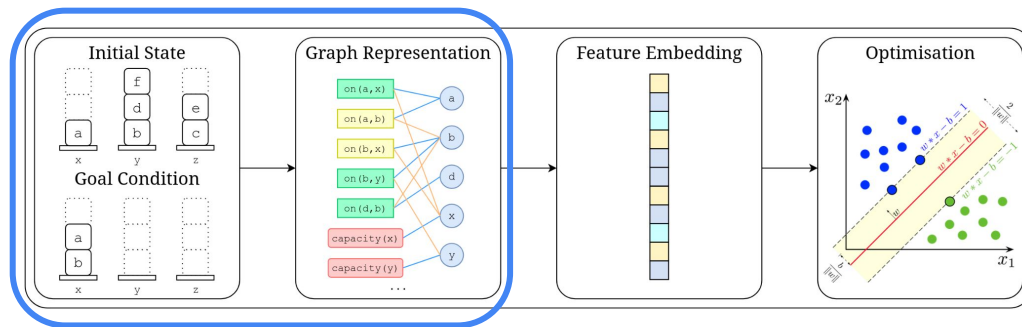


```
pip install wlplan
```



<https://dillonzchen.github.io/wlplan/>

Graph representations of planning tasks



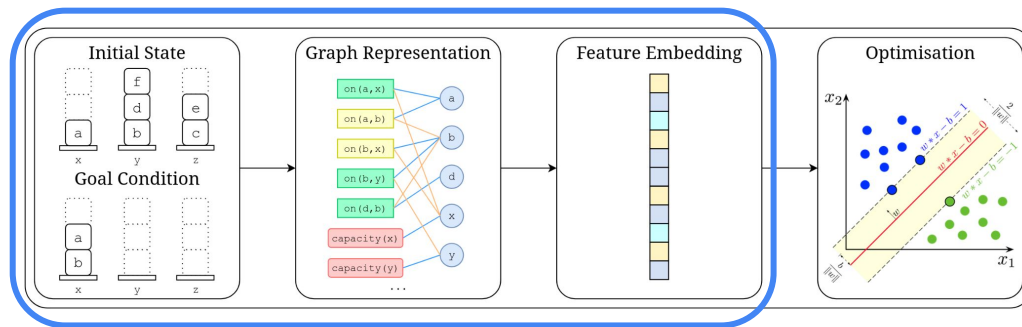
Python

```
g_generator = init_graph_generator(  
    graph_representation="ilg", domain=domain  
)  
graphs = g_generator.to_graphs(dataset)
```

C++

```
graph_generator: :ILGGenerator g_generator =  
    graph_generator: :ILGGenerator(domain);  
std::vector<graph_generator: :Graph> graphs =  
    g_generator.to_graphs(dataset);
```

Embeddings of planning tasks and graphs



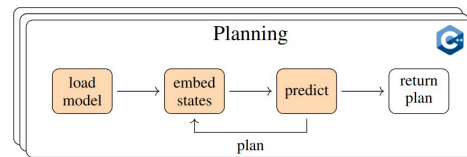
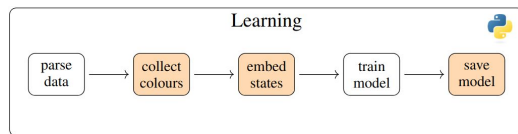
Python

```
f_generator = init_feature_generator(  
    feature_algorithm="wl", domain=domain  
)  
f_generator.collect(dataset)  
X = f_generator.embed(dataset)
```

C++

```
feature_generator: WLFeatures f_generator =  
    feature_generator: WLFeatures(domain);  
...  
std::vector<int> x =  
    f_generator.embed_state(state);
```

Serialisation of models



Python

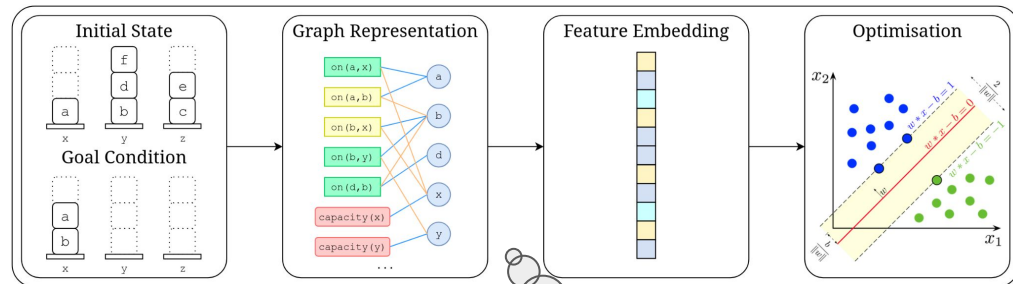
```
f_generator.save(  
    filename=x,  
    weights=predictor.get_weights(),  
)
```

C++

```
std::shared_ptr<feature_generator::Features>  
f_generator = load_feature_generator(model_file);
```

WLPlan Methodology and Configurations

WLPlan Methodology



Weisfeiler-Leman
Graph kernel

Algorithm 1: WL algorithm

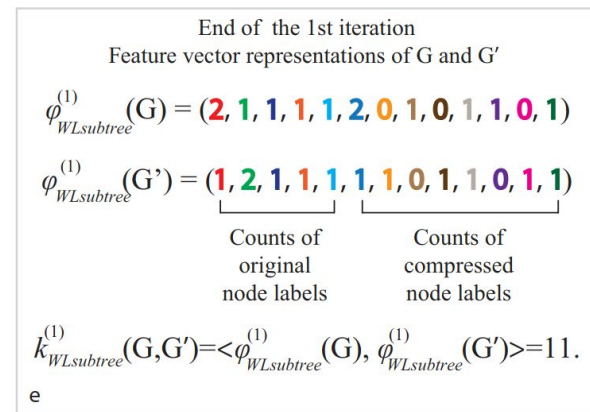
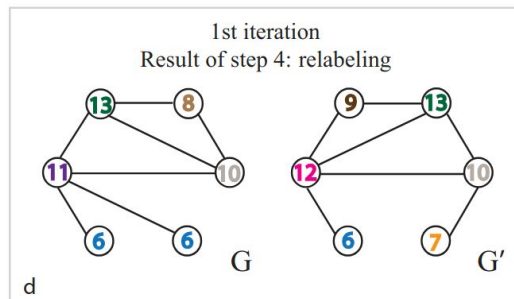
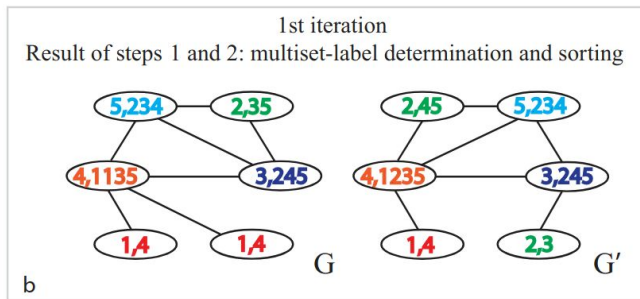
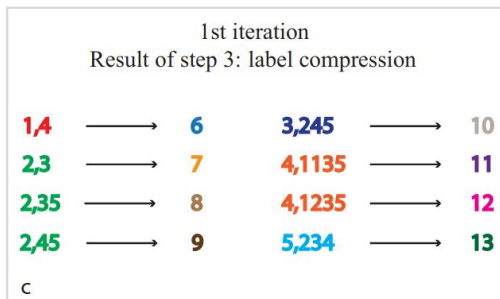
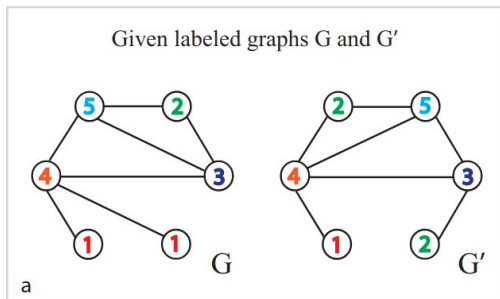
Input: A graph $G = \langle V, E, F, L \rangle$, injective HASH function, and number of iterations L .

Output: Multiset of colours.

- 1 $c^0(v) \leftarrow F(v), \forall v \in V$
- 2 **for** $l = 1, \dots, L$ **do for** $v \in V$ **do**
- 3 $c^l(v) \leftarrow$
 $\text{HASH} \left(c^{l-1}(v), \bigcup_{u \in \Sigma_E} \{ \{ c^{l-1}(u), \iota \} \mid u \in N_\iota(v) \} \right)$
- 4 **return** $\bigcup_{l=0, \dots, L} \{ c^l(v) \mid v \in V \}$

Iteratively refine node colours using
neighbour information

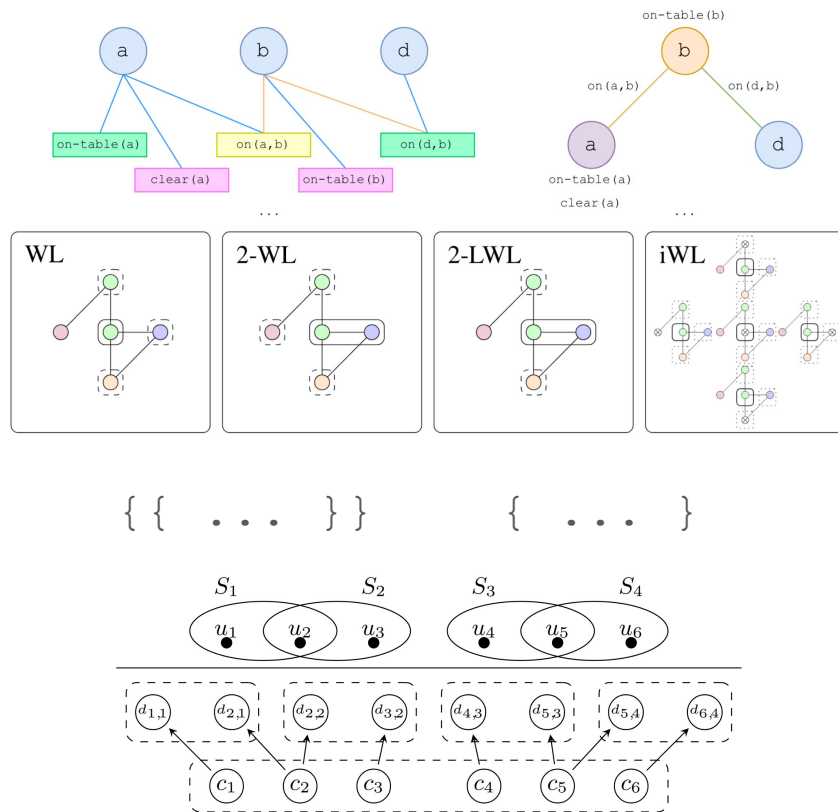
Weisfeiler-Leman Graph Kernel



WLPlan Configurations

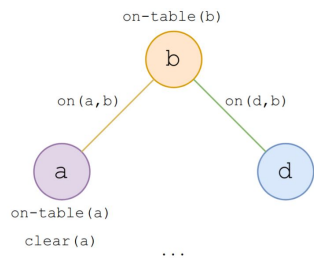
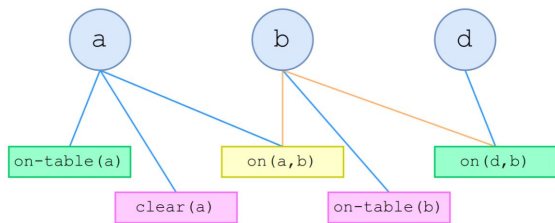
Configurations trade off between expressivity, memory and speed

Graph Representation	ILG
	PLOIG
Feature Generator	WL
	iWL
	2-LWL
	2-KWL
Hash Function	Multiset
	Set
Feature Pruner	None
	MaxSat
	iterMaxSat
	iterMaxSat+freq
Iterations	1, 2, ...



Graph Generator — Representation Options

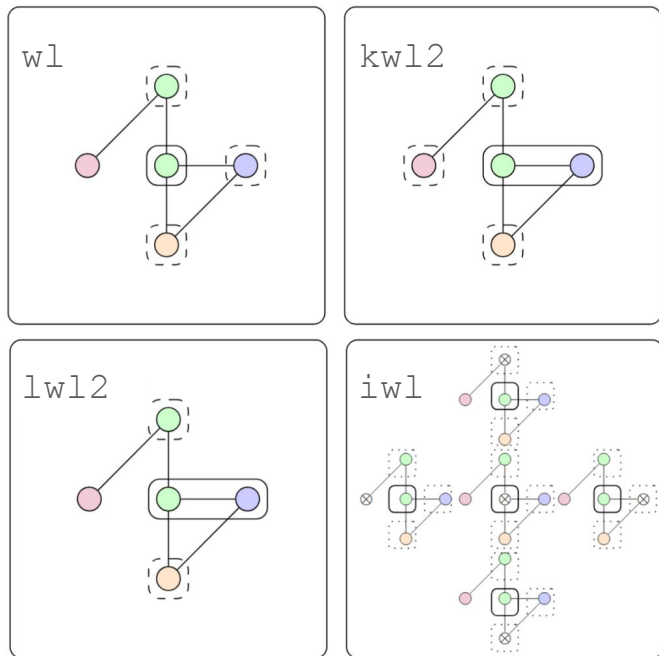
Graph representations of planning problems



```
g_generator = init_graph_generator(  
    domain : Domain,  
    graph_representation : str,  
)
```

Feature Generator — Algorithm Options

WL or an extension for generating features



```
f_generator = init_feature_generator(  
    domain : Domain,  
    feature_algorithm : str,  
    multiset_hash : bool,  
    pruning : str,  
    iterations : int,  
)
```

Feature Generator — Hash Function Options

Collect neighbour colours via multiset vs. set

Algorithm 1: WL algorithm

Input: A graph $\mathbf{G} = \langle \mathbf{V}, \mathbf{E}, \mathbf{F}, \mathbf{L} \rangle$, injective HASH function, and number of iterations L .

Output: Multiset of colours.

- 1 $c^0(v) \leftarrow \mathbf{F}(v), \forall v \in \mathbf{V}$
- 2 **for** $l = 1, \dots, L$ **do for** $v \in \mathbf{V}$ **do**
- 3 $c^l(v) \leftarrow$

$$\text{HASH} \left(c^{l-1}(v), \bigcup_{\iota \in \Sigma_{\mathbf{E}}} \{ \{ c^{l-1}(u), \iota \} \mid u \in \mathbf{N}_{\iota}(v) \} \right)$$

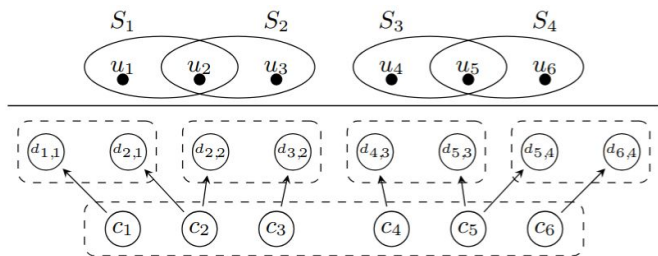
$$\text{HASH} \left(c^{l-1}(v), \bigcup_{\iota \in \Sigma_{\mathbf{E}}} \{ c^{l-1}(u), \iota \} \mid u \in \mathbf{N}_{\iota}(v) \} \right)$$

- 4 **return** $\bigcup_{l=0, \dots, L} \{ c^l(v) \mid v \in \mathbf{V} \}$
-

```
f_generator = init_feature_generator(  
    domain : Domain,  
    feature_algorithm : str,  
    multiset_hash : bool,  
    pruning : str,  
    iterations : int,  
)
```

Feature Generator — Pruning Options

Option to prune seemingly equivalent features



```
f_generator = init_feature_generator(  
    domain : Domain,  
    feature_algorithm : str,  
    multiset_hash : bool,  
    pruning : str,  
    iterations : int,  
)
```

Feature Generator — Iteration Options

Number of iterations to perform in WL

Algorithm 1: WL algorithm

Input: A graph $\mathbf{G} = \langle \mathbf{V}, \mathbf{E}, \mathbf{F}, \mathbf{L} \rangle$, injective HASH function, and number of iterations L .

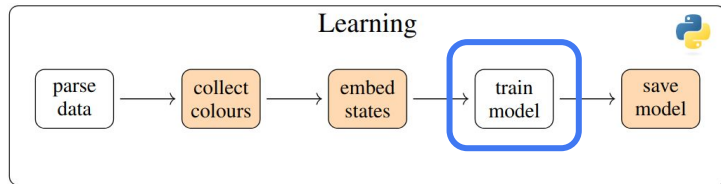
Output: Multiset of colours.

```
1  $c^0(v) \leftarrow \mathbf{F}(v), \forall v \in \mathbf{V}$ 
2 for  $l = 1, \dots, L$  do for  $v \in \mathbf{V}$  do
3    $c^l(v) \leftarrow$ 
     HASH  $\left( c^{l-1}(v), \bigcup_{\iota \in \Sigma_E} \{ \{ c^{l-1}(u), \iota \} \mid u \in \mathbf{N}_\iota(v) \} \right)$ 
4 return  $\bigcup_{l=0, \dots, L} \{ \{ c^l(v) \mid v \in \mathbf{V} \} \}$ 
```

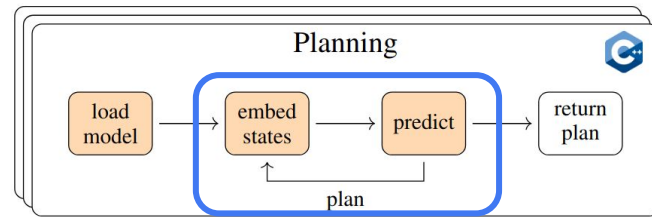
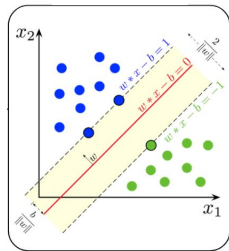
```
f_generator = init_feature_generator(
    domain : Domain,
    feature_algorithm : str,
    multiset_hash : bool,
    pruning : str,
    iterations : int,
)
```


WLPlan Use Cases

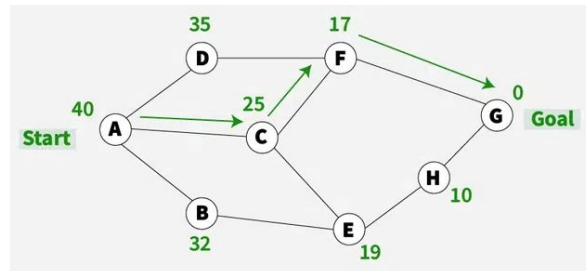
Learning Heuristics



Learn a heuristic



Employ the heuristic in search, e.g. Greedy Best First Search



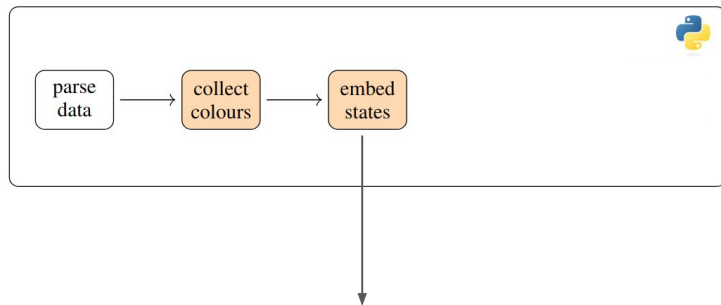
https://dillonzchen.github.io/wlplan/tutorials/3_gnns.html

[View page source](#)

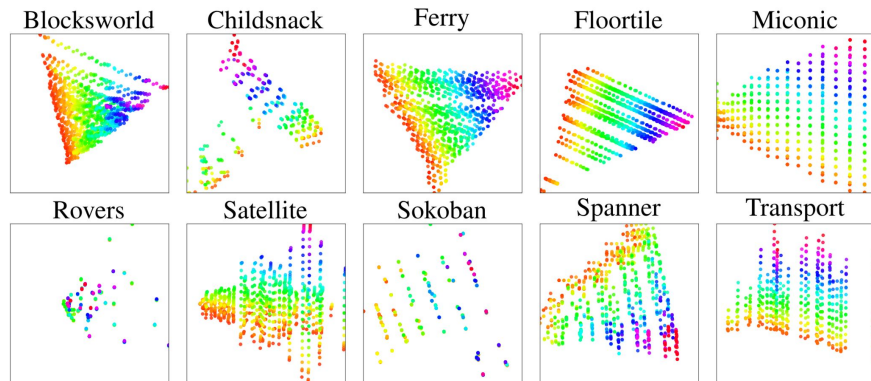
```
val_loader = DataLoader(dataset_loader.get_loader('val'), batch_size=batch_size,
                        num_workers=num_workers, shuffle=False)

return train_loader, val_loader
```

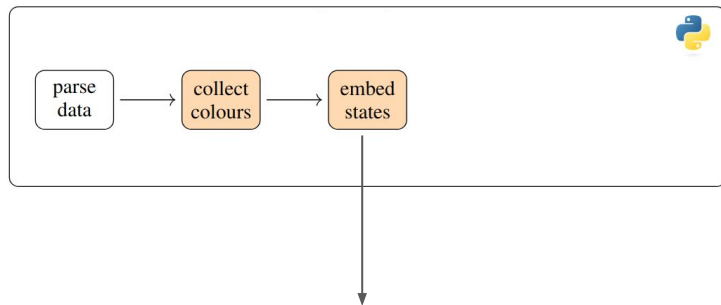
Visualising Planning Domains



perform dimensionality Reduction
(e.g. PCA, t-SNE)



Expressivity Testing



Symmetries and Expressive Requirements for Learning General Policies

Dominik Drexler¹, Simon Ståhlberg², Blai Bonet³, Hector Geffner²

¹Linköping University, Sweden

²RWTH Aachen University, Germany

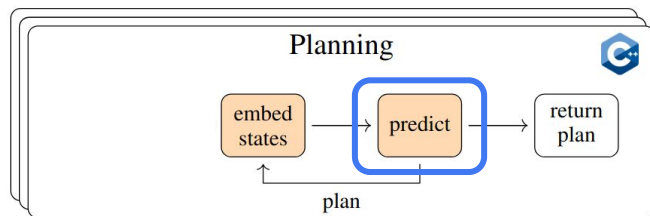
³Universitat Pompeu Fabra, Spain

dominik.drexler@liu.se, simon.stahlberg@gmail.com,
bonetblai@gmail.com, hector.geffner@ml.rwth-aachen.de

	Blocksworld	Childsnack	Ferry	Floortile	Miconic	Rovers	Satellite	Sokoban	Spanner	Transport
WL	0	0	0	22	0	0	0	129	68	0
2-LWL	0	0	0	–	0	–	–	–	14	0
iWL	0	0	0	22	0	–	0	110	68	0

Table 1: Number of pairs of indistinguishable states encoded as ILGs with respect to h^* values with WL algorithms on IPC23LT domains. – indicates that the memory limit was reached while embedding the entire dataset with the corresponding WL algorithm. Lower (↓) values are better.

Synthesising Novelty Heuristics



$$h_{\text{pn}}^{\mathcal{F}}(s) = \min_{\sigma \subseteq \mathcal{F}(s), \text{s.t. } \nexists t \in C, \sigma \subseteq \mathcal{F}(t)} |\sigma|$$

Lipovetzky and Geffner 2017 for $\mathcal{F}(s) = \text{id}(s)$

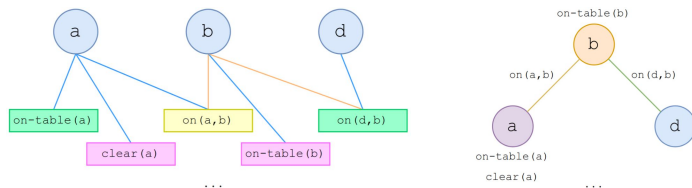
$$h_{\text{qb}}^{\mathcal{F}}(s) = \begin{cases} -|\mathcal{F}_{\text{new}}(s)|, & \text{if } |\mathcal{F}_{\text{new}}(s)| > 0, \\ +|\mathcal{F}_{\text{old}}(s)|, & \text{otherwise,} \end{cases}$$

$$\mathcal{F}_{\alpha}(s) = \{p \in \mathcal{F}(s) \mid h(s) \triangleq \min_{t \in C, p \in \mathcal{F}(t)} h(t)\}$$

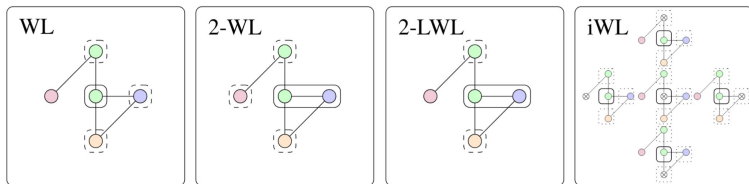
Katz et al. 2017 for $\mathcal{F}(s) = \text{id}(s)$

	at	wl	at;wl	at wl
qb	508	462	527	543
pn	612	556	599	644

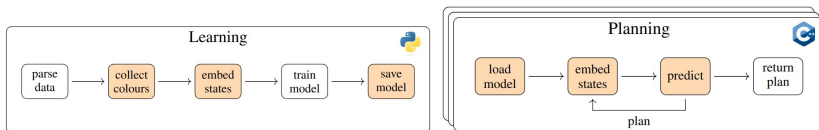
1. Graph representations of planning tasks



2. Embeddings of planning tasks and graphs



3. Serialisation of models





WLPlan

GET STARTED

Installation

Introduction by Example

TUTORIALS

Graphs and NetworkX

Graph Neural Networks

PACKAGE REFERENCE

wlplan

wlplan.data

wlplan.feature_generator

wlplan.graph_generator

wlplan.planning

WLPlan Documentation

Edit on GitHub

WLPlan Documentation

WLPlan is a library for generating graph representations and feature embeddings of PDDL planning problems and states for machine learning tasks. WLPlan currently supports both classical and numeric planning problems.

The main pipeline in WLPlan consists of (1) converting planning problems and states into graphs, and (2) synthesising feature embeddings by running a variant of the Weisfeiler-Leman (WL) algorithm on the resulting graph.

Initial State

Convert to Graph

Generate Features

The entire pipeline of generating feature embeddings from PDDL planning problems can be performed in a matter of steps with WLPlan:

```
Feature_generator = Init_Feature_generator(domain)
Feature_generator.on_init(dataset)
X = Feature_generator.embed(dataset)
```

Get Started

Installation

Introduction by Example

Tutorials

Graphs and NetworkX

Graph Neural Networks

Package Reference

wlplan

wlplan.data

wlplan.feature_generator

wlplan.graph_generator

wlplan.planning

<https://github.com/dillonzchen/wlplan>

Implementing a L4P Pipeline from Scratch

1. Parsing training data
2. Manipulating training data
3. Building a model
4. Training a model
5. Evaluating a model

<https://github.com/l4p-tutorial/heuristic-learning>