

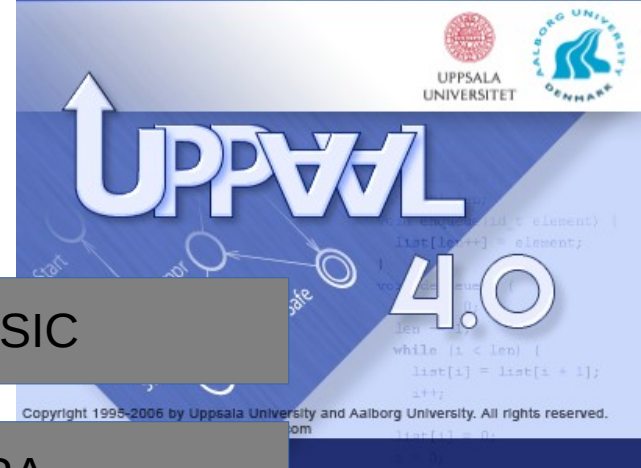


Planning using UPPAAL SMC & UPPAAL STRATEGO

Peter G. Jensen, May the Fourth, 2021



AALBORG
UNIVERSITET



Overview

Timed Automata

- Decidability (regions)
- Symbolic Verification (zones)

Priced Timed Automata

- Decidability (priced regions)
- Symbolic Verification (priced zones)

Timed Games

- Decidability (regions)
- Symbolic Verification (zones)

Stochastic Priced Timed Automata

- Stochastic semantics
- Statistical methods

Stochastic Priced Timed Game

- Stochastic Semantics
- Reinforcement Learning

Decidable/Solvable
Not expressive
Hard guarantees

Undecidable/Unsolvable
Very expressive
"Measured" guarantees

CLASSIC

CORA

TIGA

ECDAR

TRON

SMC

STRATEGO

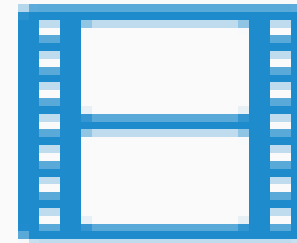
Outline

Stochastic Timed Automata (SMC)

- Markov Chains
- Stochastics and Time
- Common Modeling Mistakes
- Exercise!
- Examples

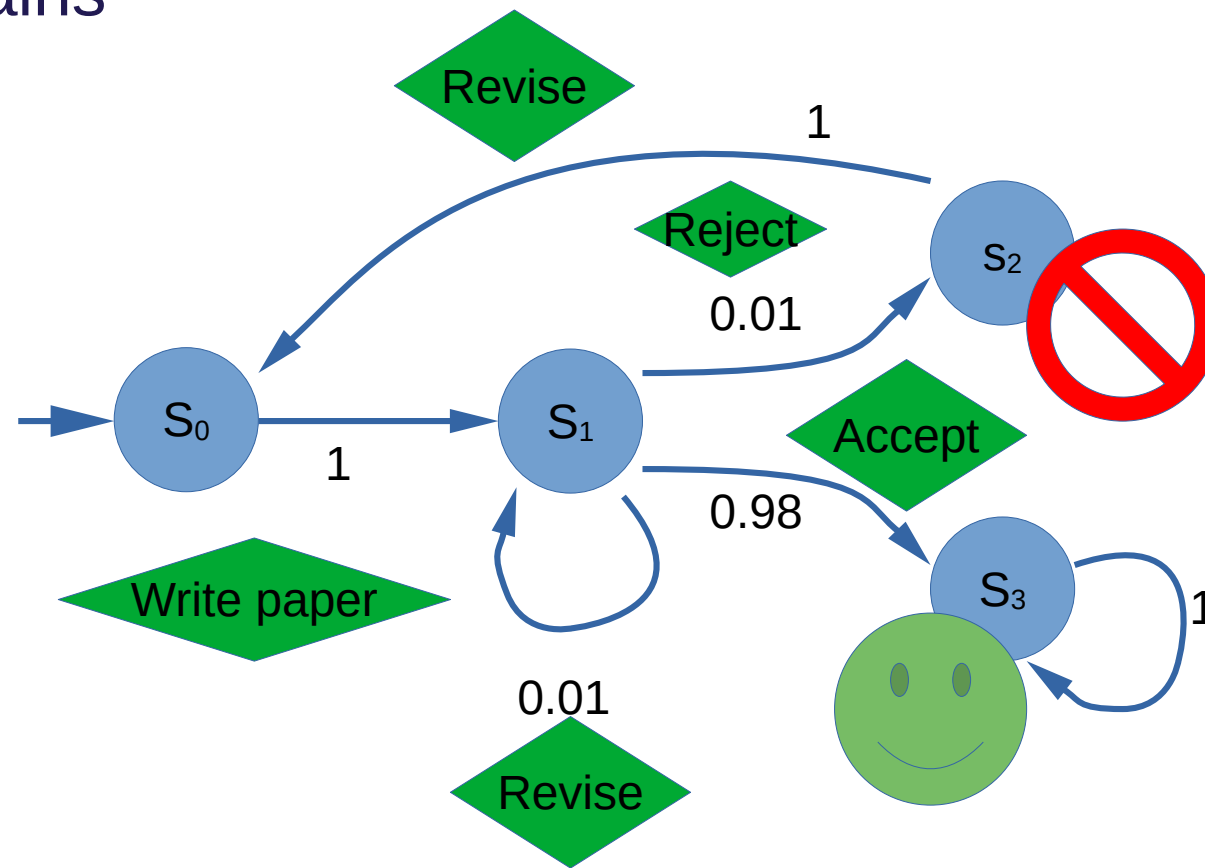
Stochastic Priced Timed Games (Stratego)

- Markov Decision Processes
- Q-learning (Tabulation)
- Euclidean Markov Decision Processes
- Stratego
 - Safe learning!
 - Partial observability
- Exercise!
- Examples



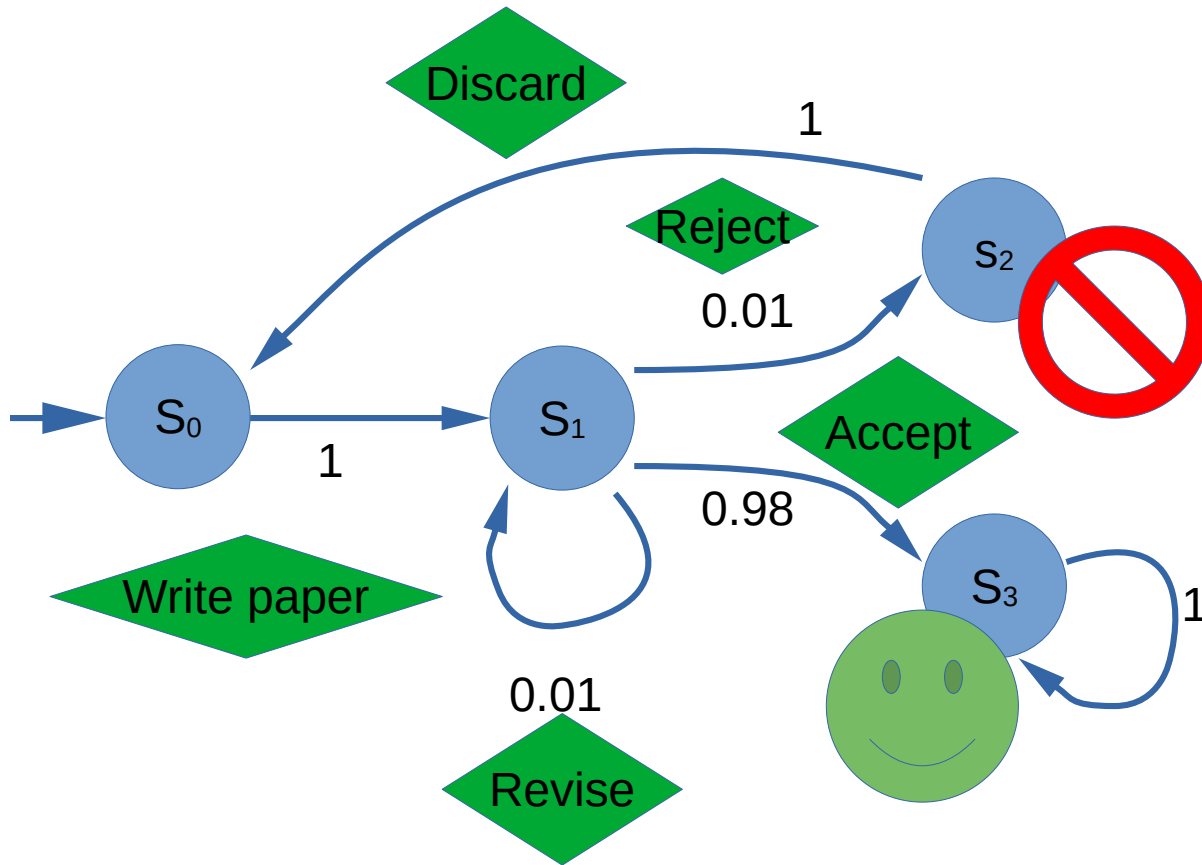
Stochastic Semantics

Markov Chains



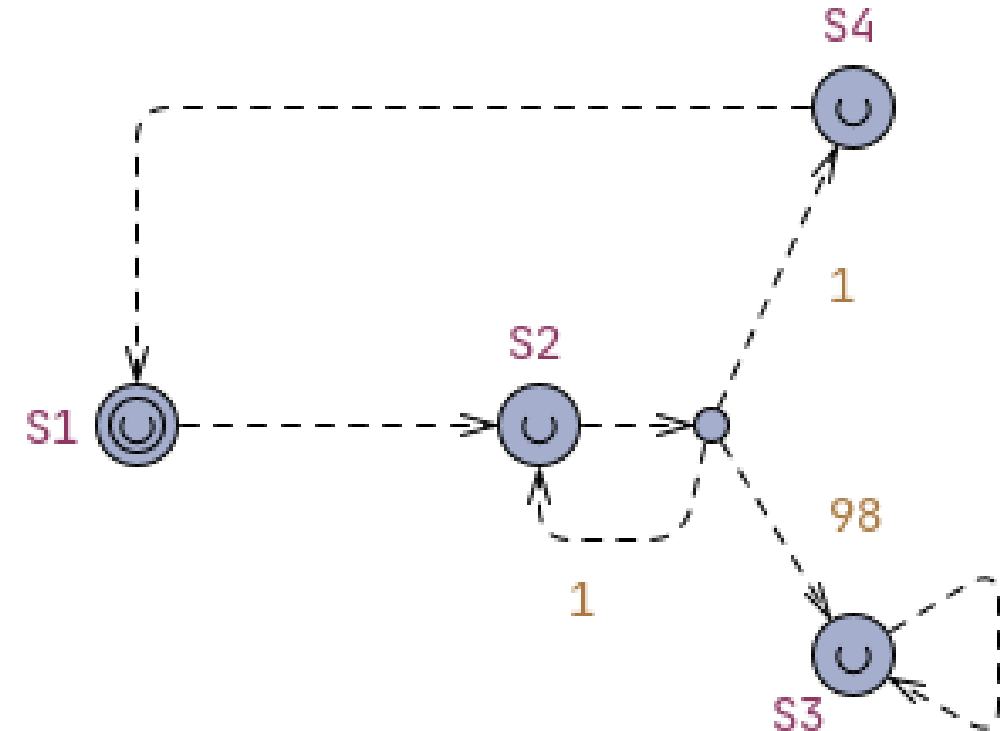
$$\Pr([] \vdash S_2) = 0.98/0.99 \approx 0.98989898\dots$$

Stochastic Semantics Markov Chains



$$\Pr([] \vdash S_2) = 0.98/0.99 \approx 0.98989898\dots$$

... In UPPAAL!



Notice, no time yet, all locations are urgent!

Stochastic Semantics vs Nondeterministic Semantics

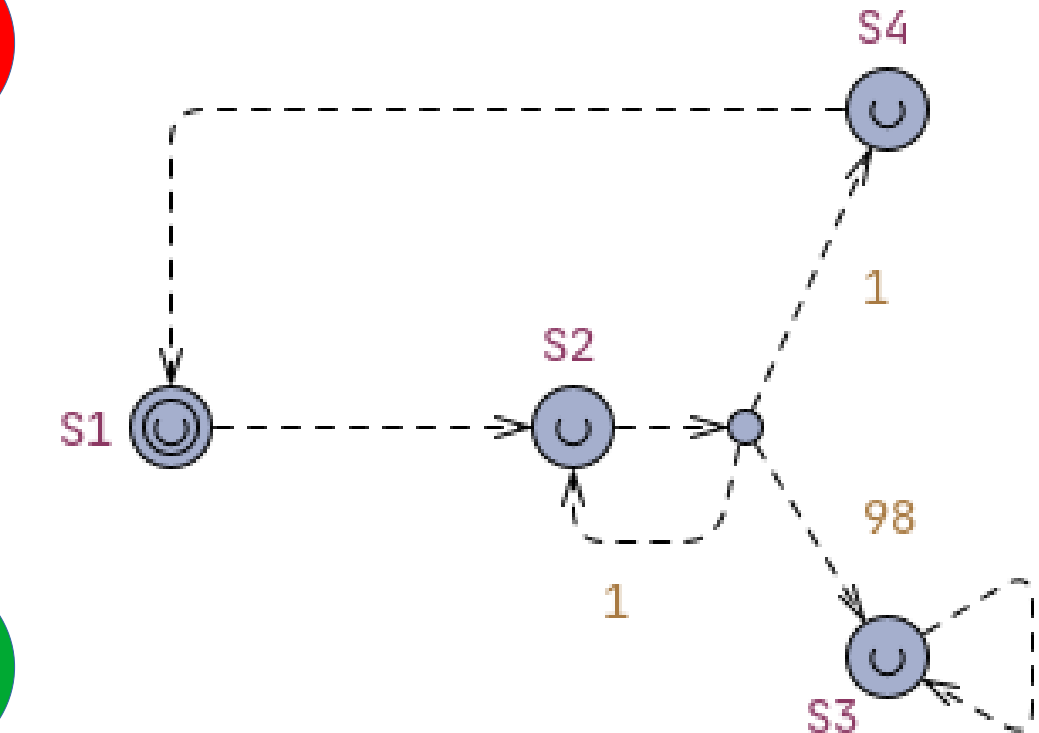
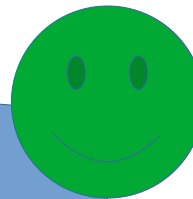
We can loop
 $S1, S2^*$ forever

Nondeterministic:
 $A \langle \rangle S3$



Will we always reach $S3$?

Stochastic/probabilistic:
 $\Pr(\langle \rangle S3)$



Probability tends to 1
over infinite execution

Stochastic Semantics vs Nondeterministic Semantics

We can loop
S1, S2* forever

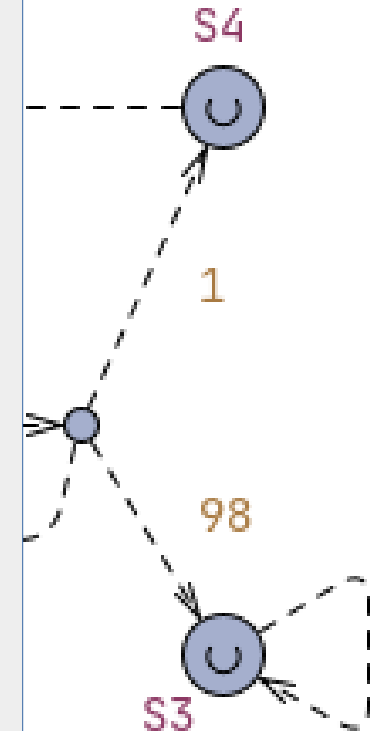
Nondet
A<

Will we alwa

Stochastic
Pr(<> S3)

Lesson:
 $[\text{Pr}(X) = 1] \neq [A \langle \rangle X]$

Probability tends to 1
over infinite execution



Mini Exercise

You want to play a game of dices – but you only have a coin available (heads/tails). Implement an emulator of a 6-dice using coinflips.

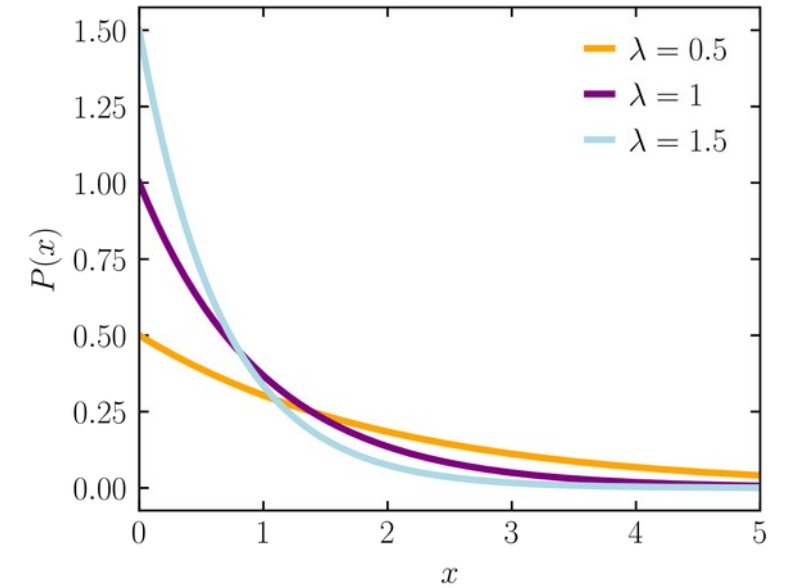
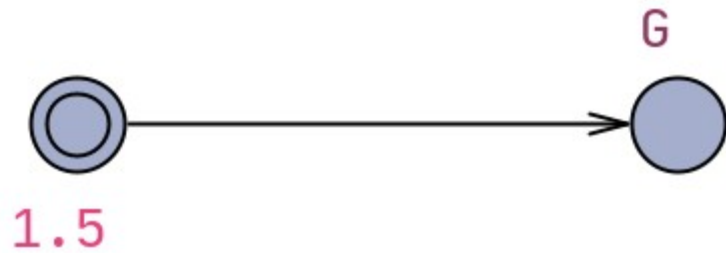
Ensure that the dice is fair!
(What is a good specification?)

You can use

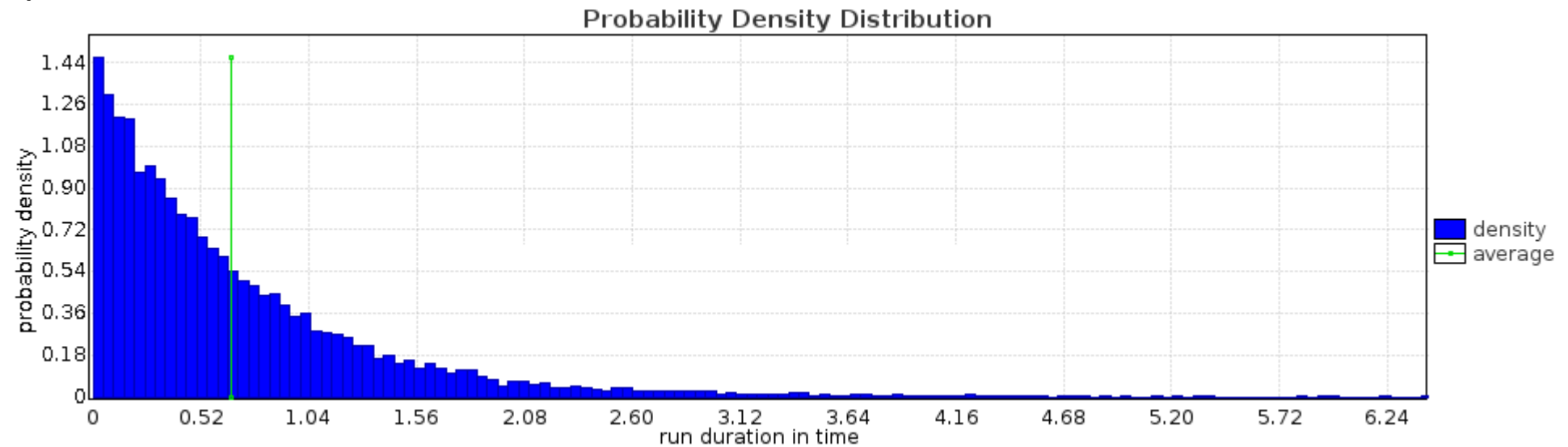
$\text{Pr}[\leq 1](\langle \rangle \text{ Process.L})$
to estimate the probability of ending in a location L of Process

Tip:
Make all (but terminal) locations urgent!

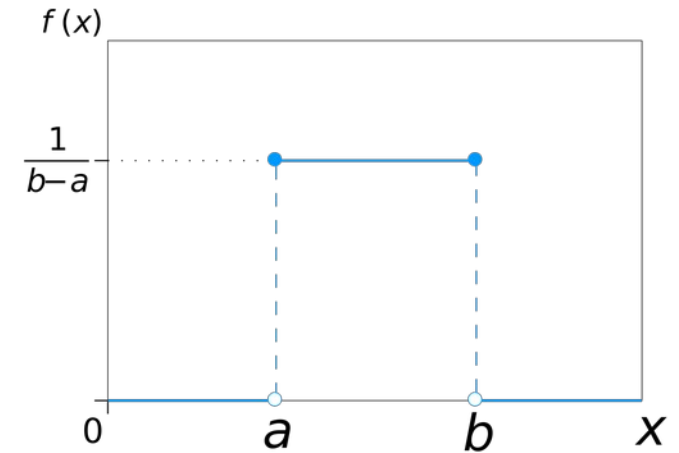
Stochastic Semantics and Time



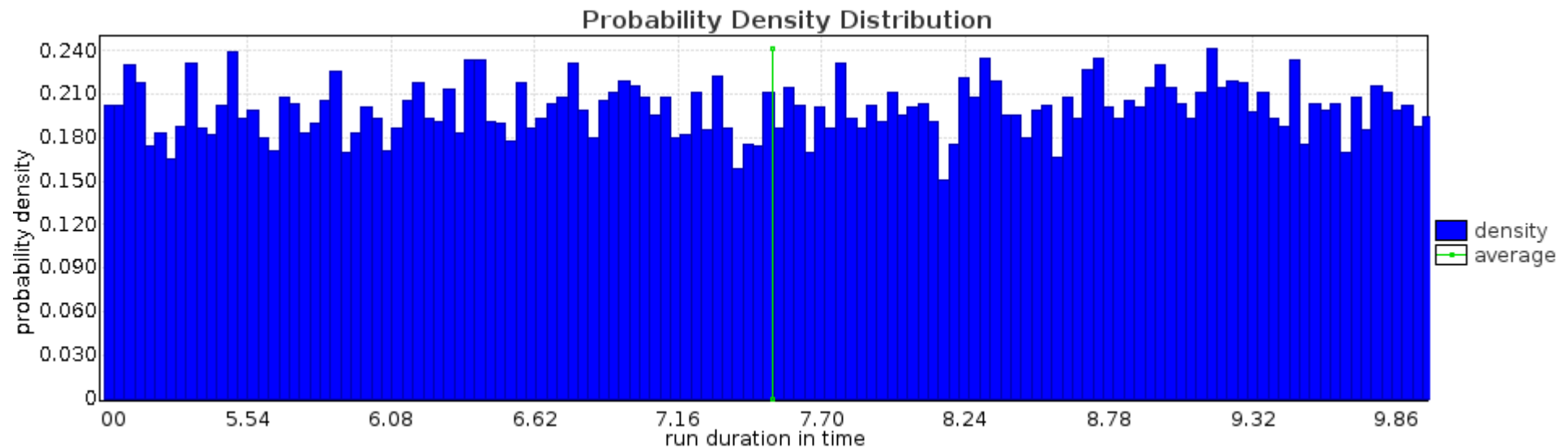
$\Pr(< > G) =$



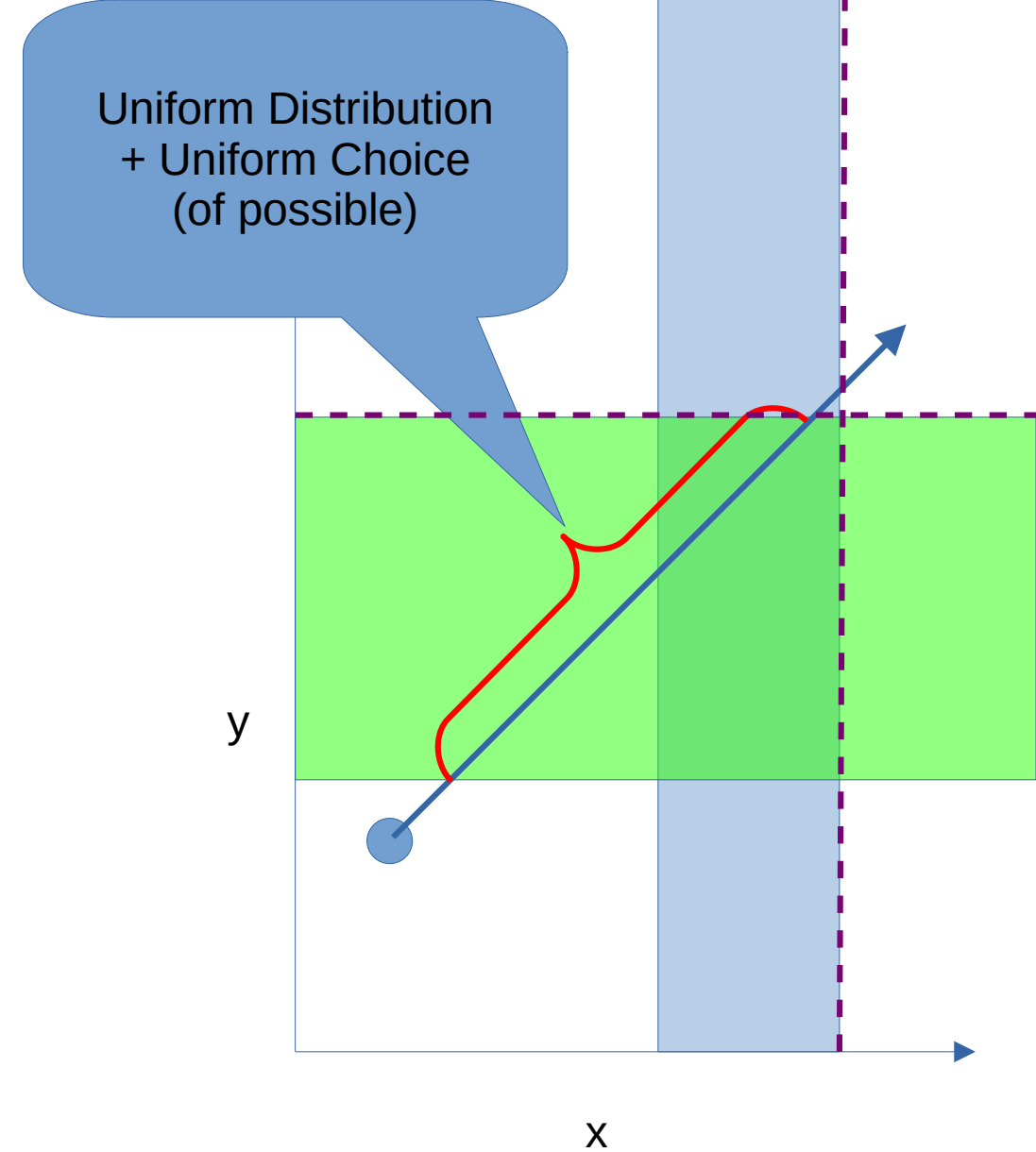
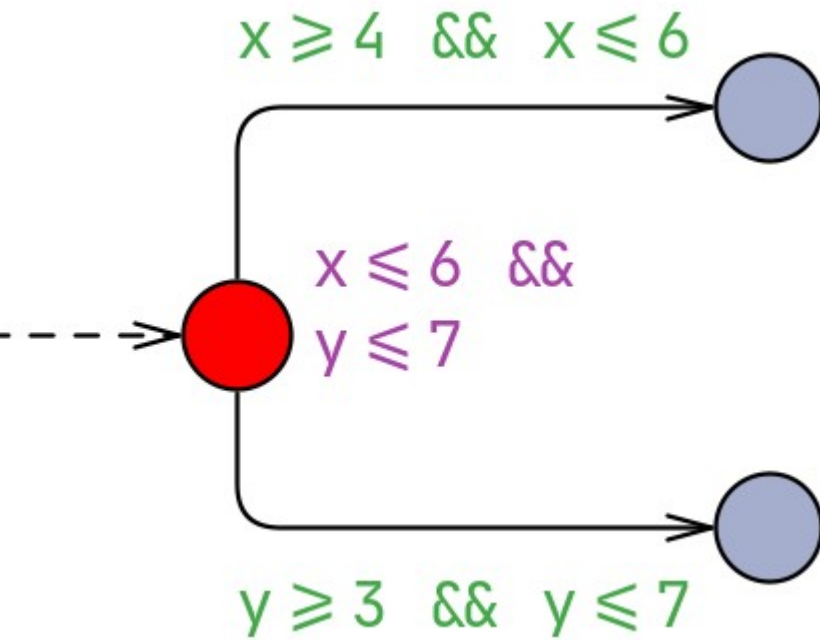
Stochastic Semantics and Time



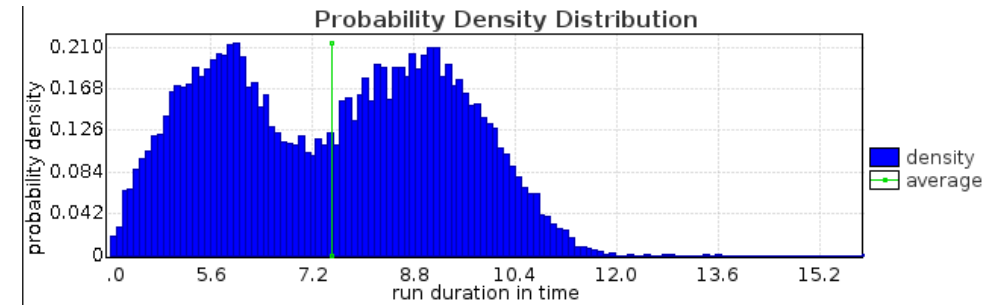
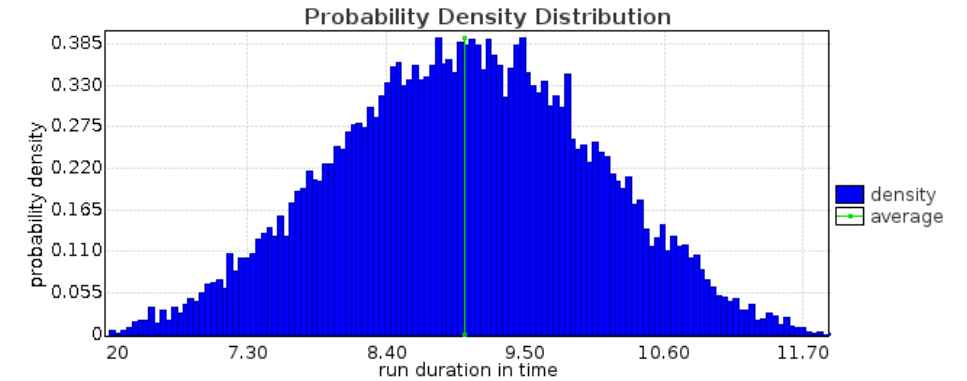
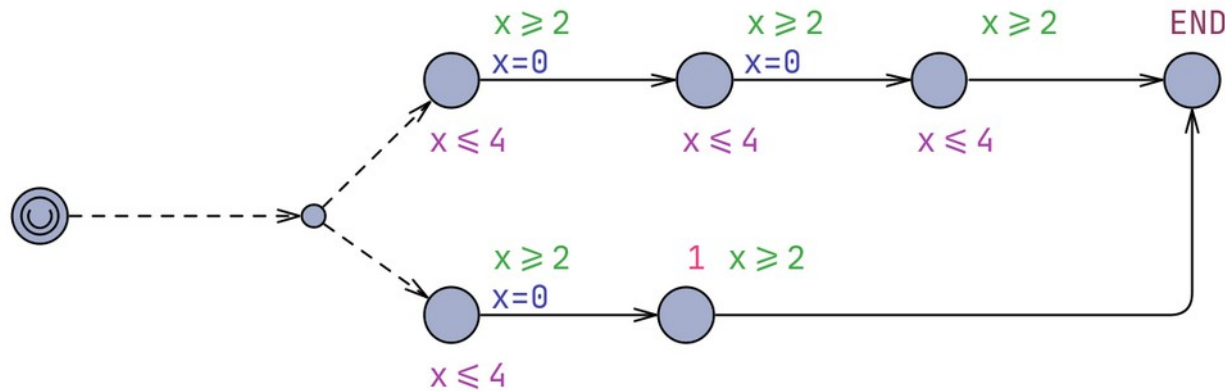
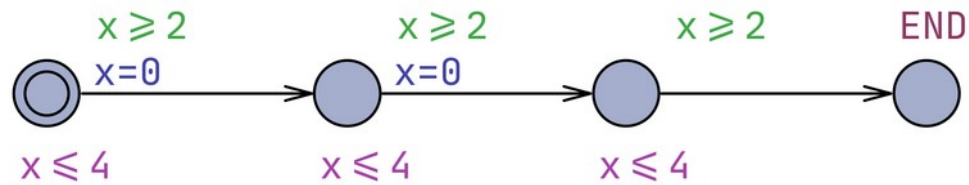
$\Pr(\langle \rangle G) =$



Stochastic Semantics and Multiple Temporal Choices



Stochastic Semantics and Phase-Type Distributions

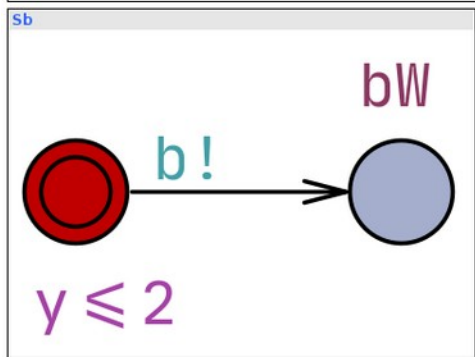
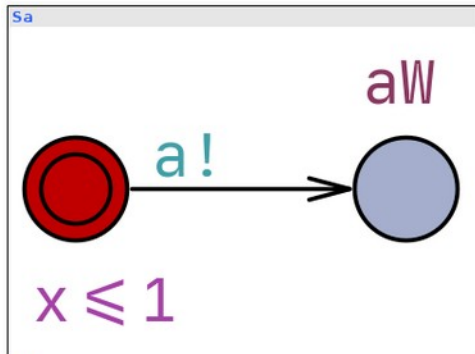
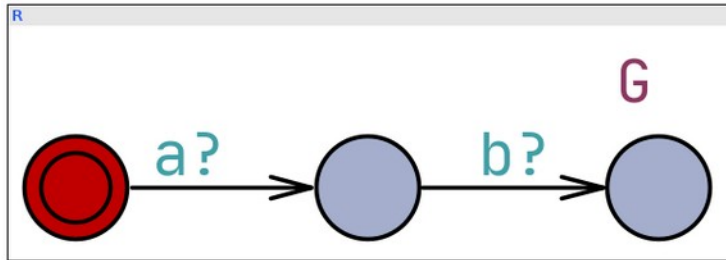


Any distribution can be encoded
to arbitrary precision



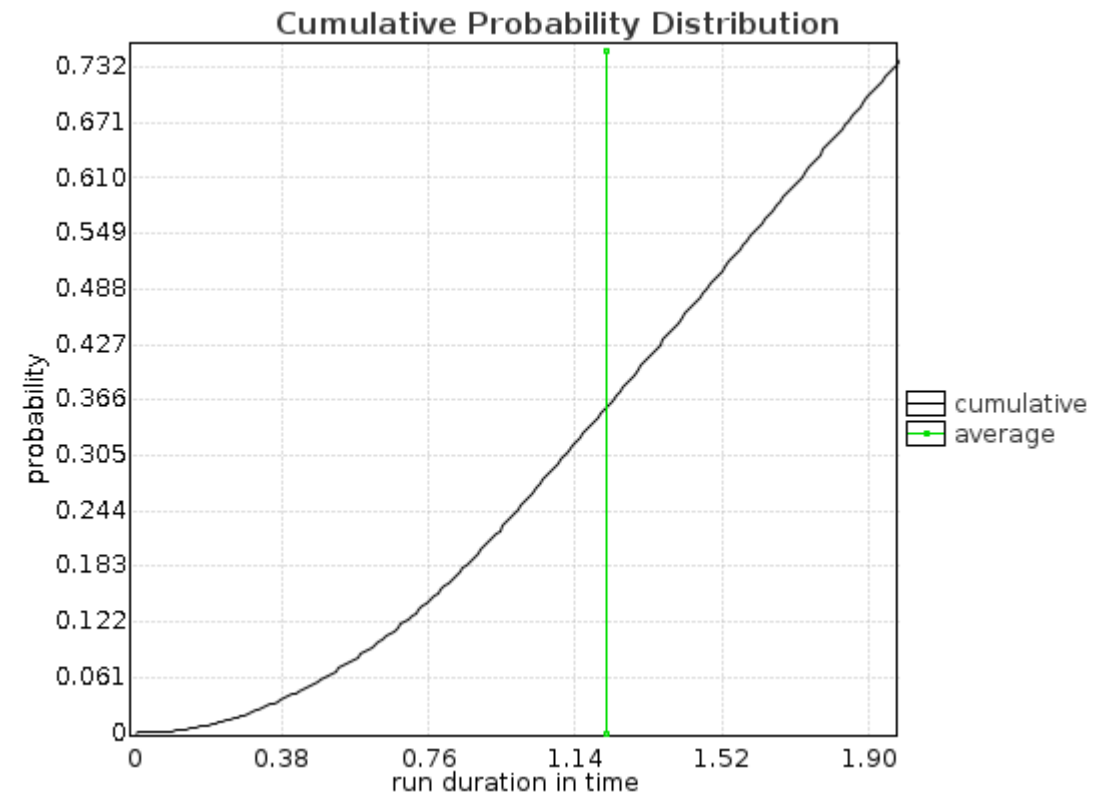
Calls to external C-function allow
for arbitrary distributions

Stochastic Semantics and Multiple Automata



Race Between components

$$\Pr(<> G) = 0.75$$



Queries of SMC

Evaluation

```
Pr[<=100] (<> expr)  
Pr[<=100] ([] expr)
```

Hypothesis Testing

```
Pr[<=100] (<> expr) >= 0.1  
Pr[<=100] ([] expr) >= 0.1
```

Comparison

```
Pr[<=100] (<> expr) >= Pr[<=10] (<> expr2)
```

Expected Value

```
E[<=100;100] (min: expr)  
E[<=100;100] (max: expr)
```

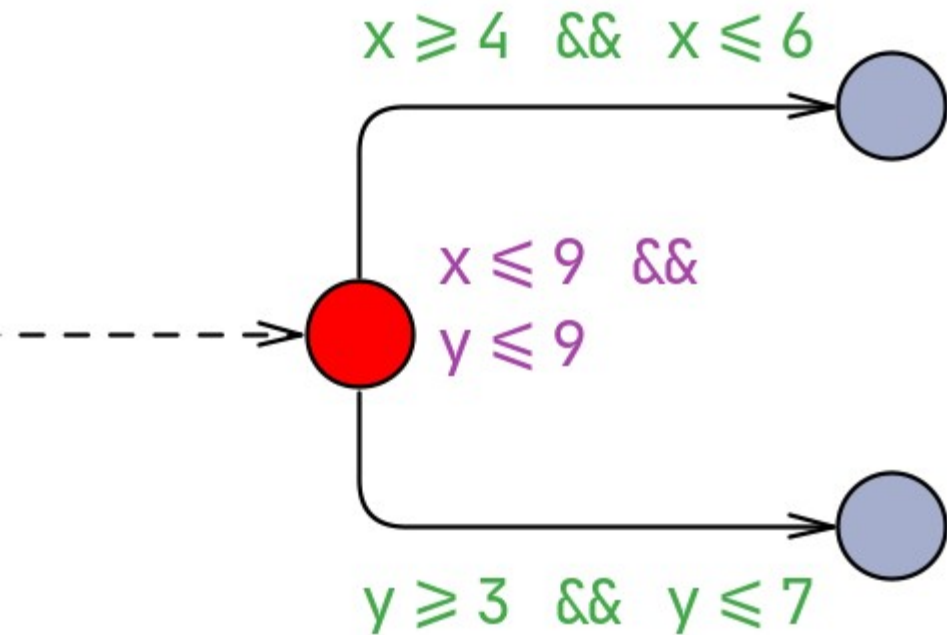
Simulation

```
simulate 1 [<=100] {e1, e2}
```

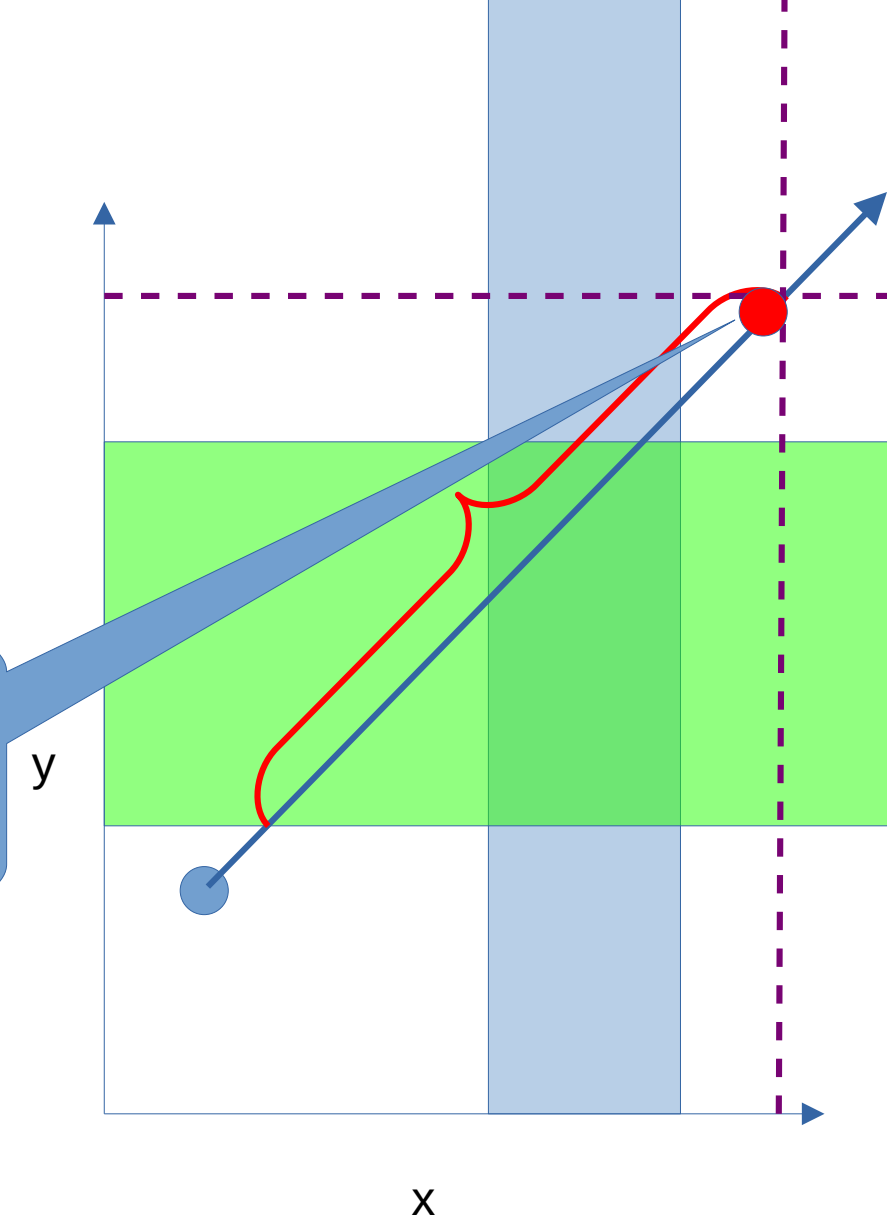
Notice:
Timebound is always needed!
It is a sampling-based method.

Common Modelling Mistakes

Time Locks



Time cannot progress!
No escape!



Common Modelling Mistakes

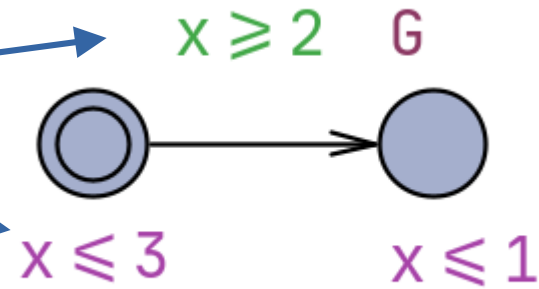
Model Sanity

Delay must be
In $[2,3]$

$\Leftrightarrow$

x is in $[2,3]$ after delay

Different to Model Checking (Classic)
We cannot impose constraints backwards



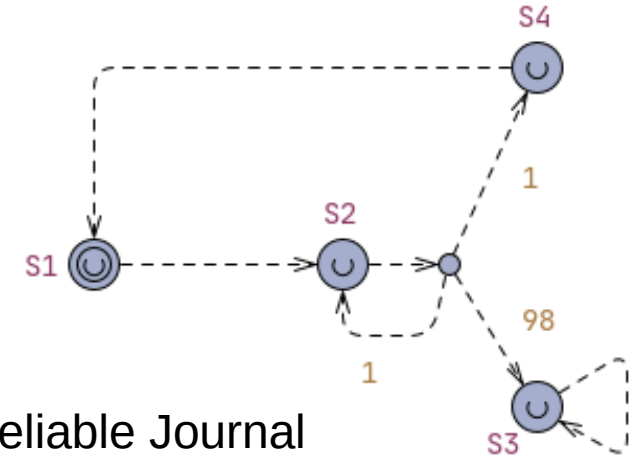
Valid successor must have
 $x \leq 1$

Simulation allows for constructions
that make such analysis/imposition
theoretically impossible

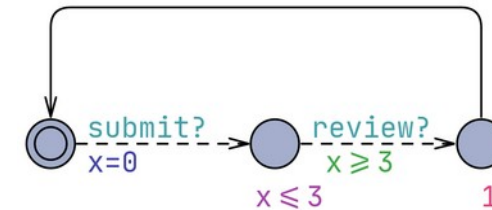
Exercises

1) Temporal Papers

- Extend the paper-writing model s.t. the researcher starts writing a paper uniformly between 1 and 6 months.
- It takes the researcher one month (with an exponential distribution) to revise the paper.
- If the paper is rejected, an exponential recuperating period is mandated (2 months).
- If the paper is accepted, the researcher starts over.
- What is the expected number of papers produced over a 5 year period?
- Connect three researchers with the “Unreliable Journal”, what is the throughput (you may ignore the review signal)?

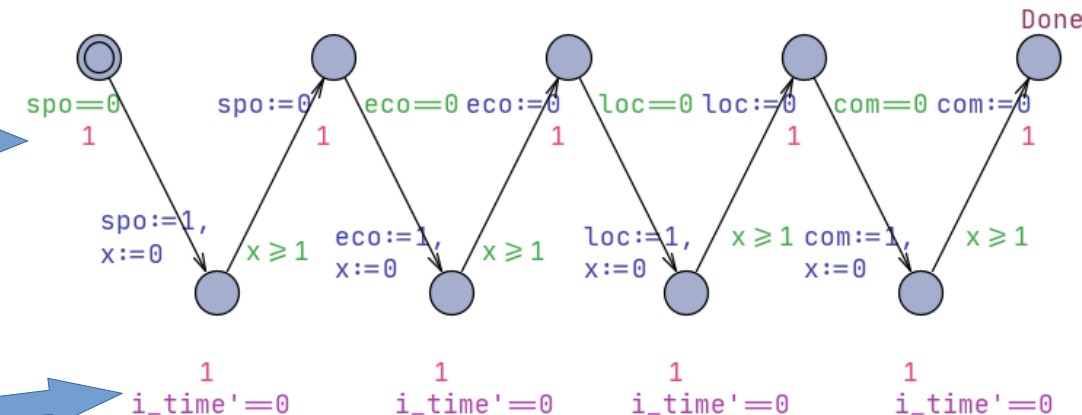


The Unreliable Journal



2) Stochastic Jobshop Scheduling

- Extend your Jobshop scheduling problem by adding distributions.
 - Assume the upperbound for reading is +5 timeunits.
 - Use a variable for the exponential rate for the “idle”-locations
- Study the effect of different exponential rates.
- Study the idle-duration of Kim
 - add a clock `i\_time` with `i\_time'==0` in the invariant in non-idle locations



Field Measurements

Generated Measures

GPS-Log

VisuAAL

RSSI-Log

Synchronized Time Series

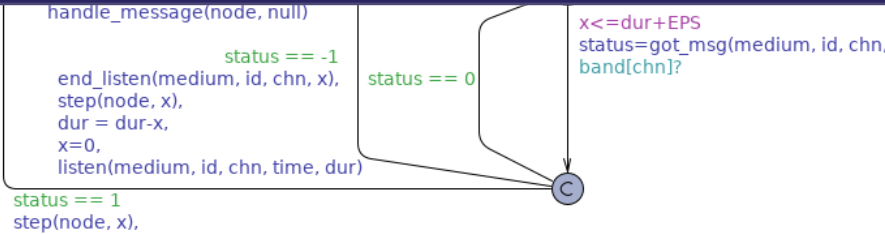
Radio Model

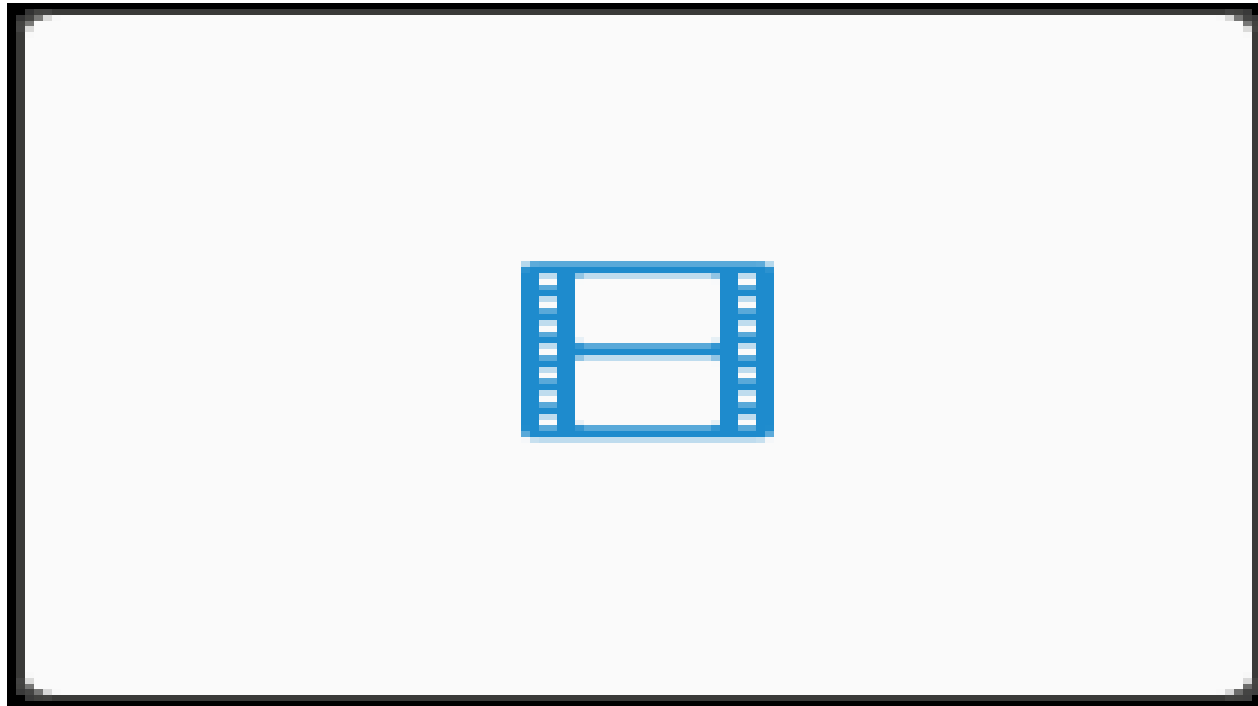
Connection Probability

UPPAAL Simulator

- Coordination
- Synchronization
- Data Collection

- Send
- Receive
- Sleep
- Inspect

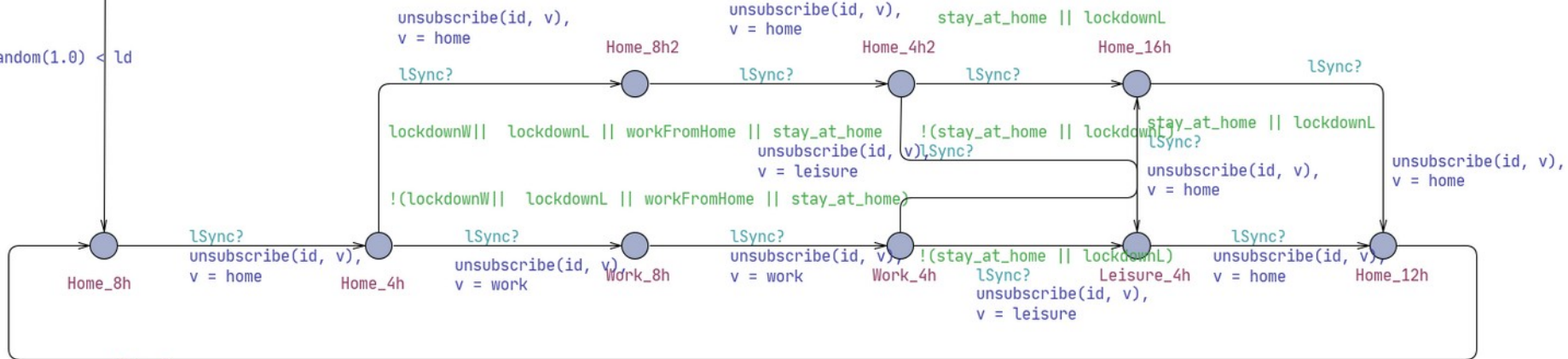




```

lSync?
v = home,
workFromHome = random(1.0) < ld

```



lSync?

```

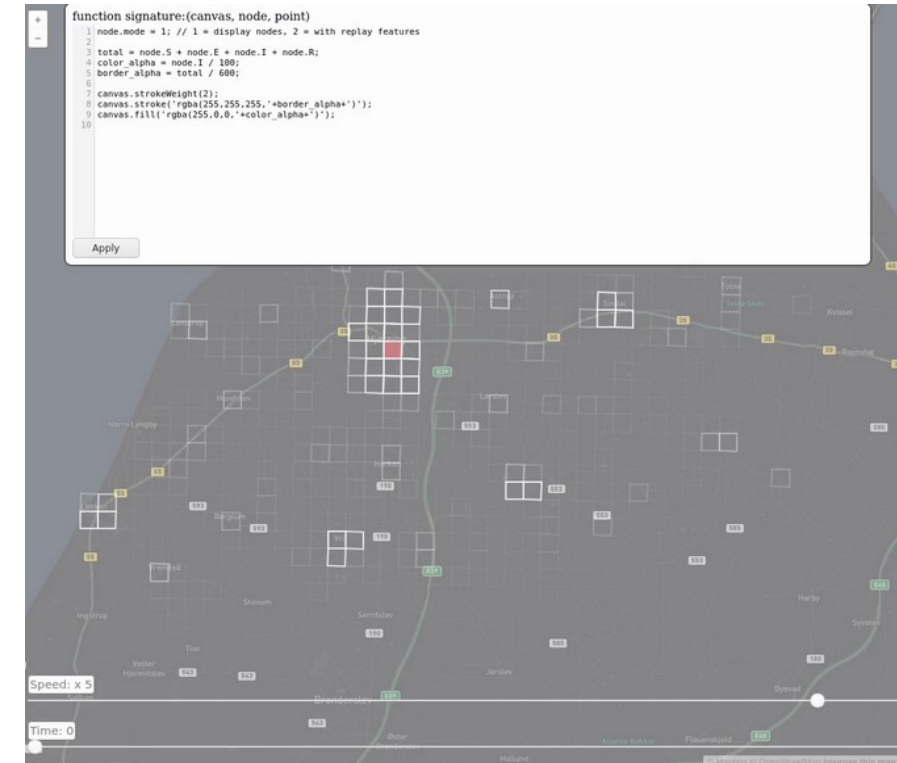
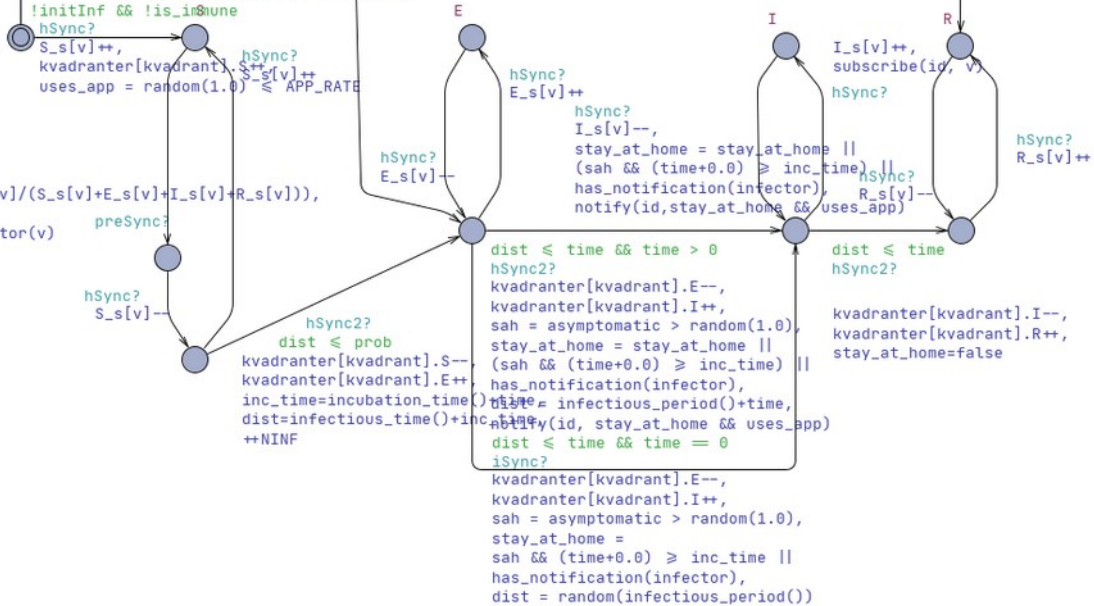
is_immune
hSync?
R_s[v]++, kvadrant[kvadrant].R++,
uses_app = random(1.0) <= APP_RATE
initInf && ! is_immune
iSync?
inc_time=incubation_time(),
dist = infectious_time()+inc_time,
prob = infectious_period(),
prob = random(dist+prob),
inc_time = inc_time - prob,
dist = dist - prob,
kvadrant[kvadrant].E++,
uses_app = random(1.0) <= APP_RATE
!initInf && ! is_immune
hSync?
S_s[v]++, kvadrant[kvadrant].S++,
uses_app = random(1.0) <= APP_RATE

```

```

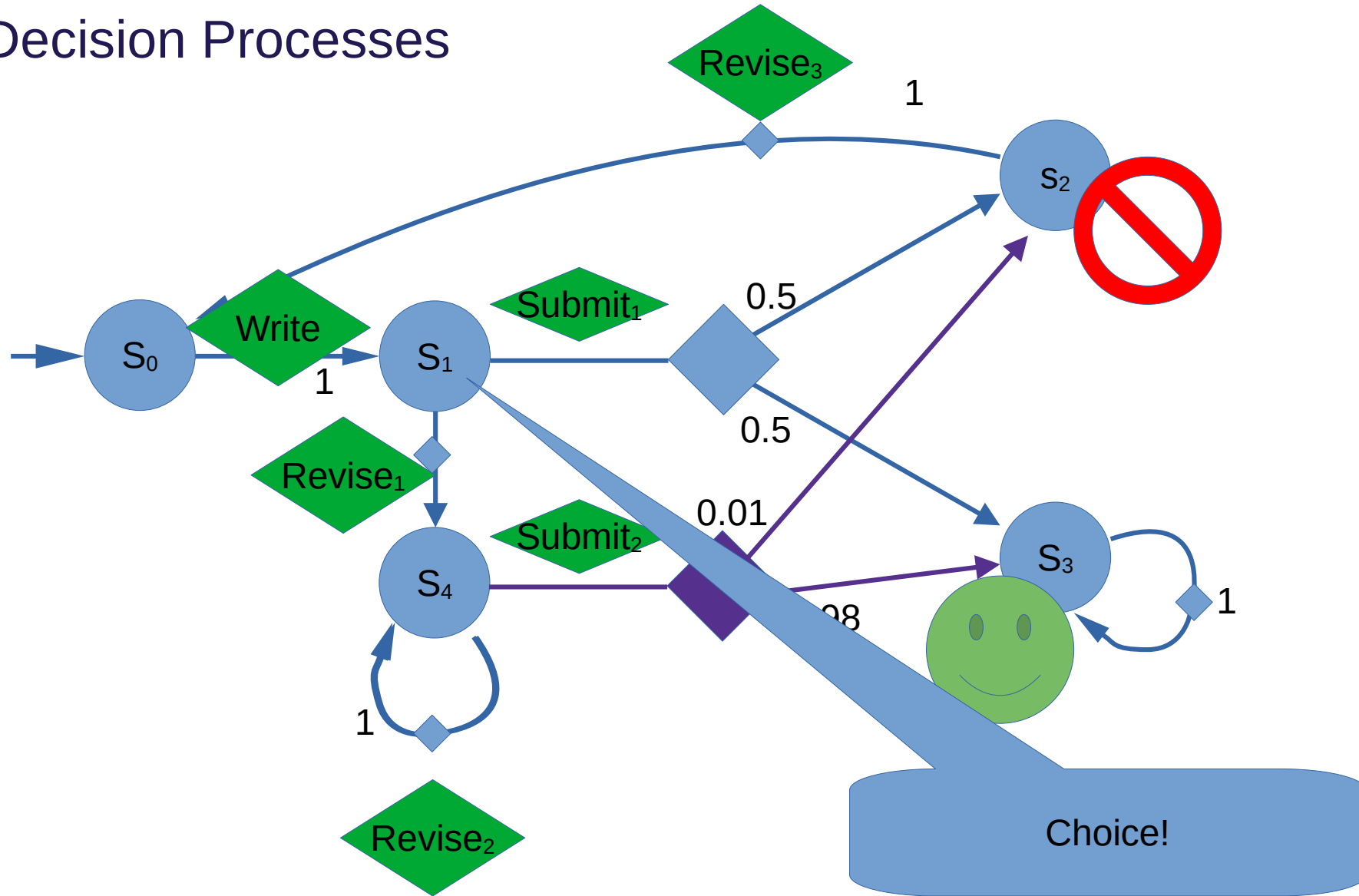
prob = cdf(beta*I_s[v]/(S_s[v]+E_s[v]+I_s[v]+R_s[v])),
dist = random(1.0),
infectior=sample_infectior(v)

```



Stochastic Optimization of Stochastic Priced Timed Games using Uppaal Stratego

Markov Decision Processes



Markov Decision Processes

We could have costs on the actions

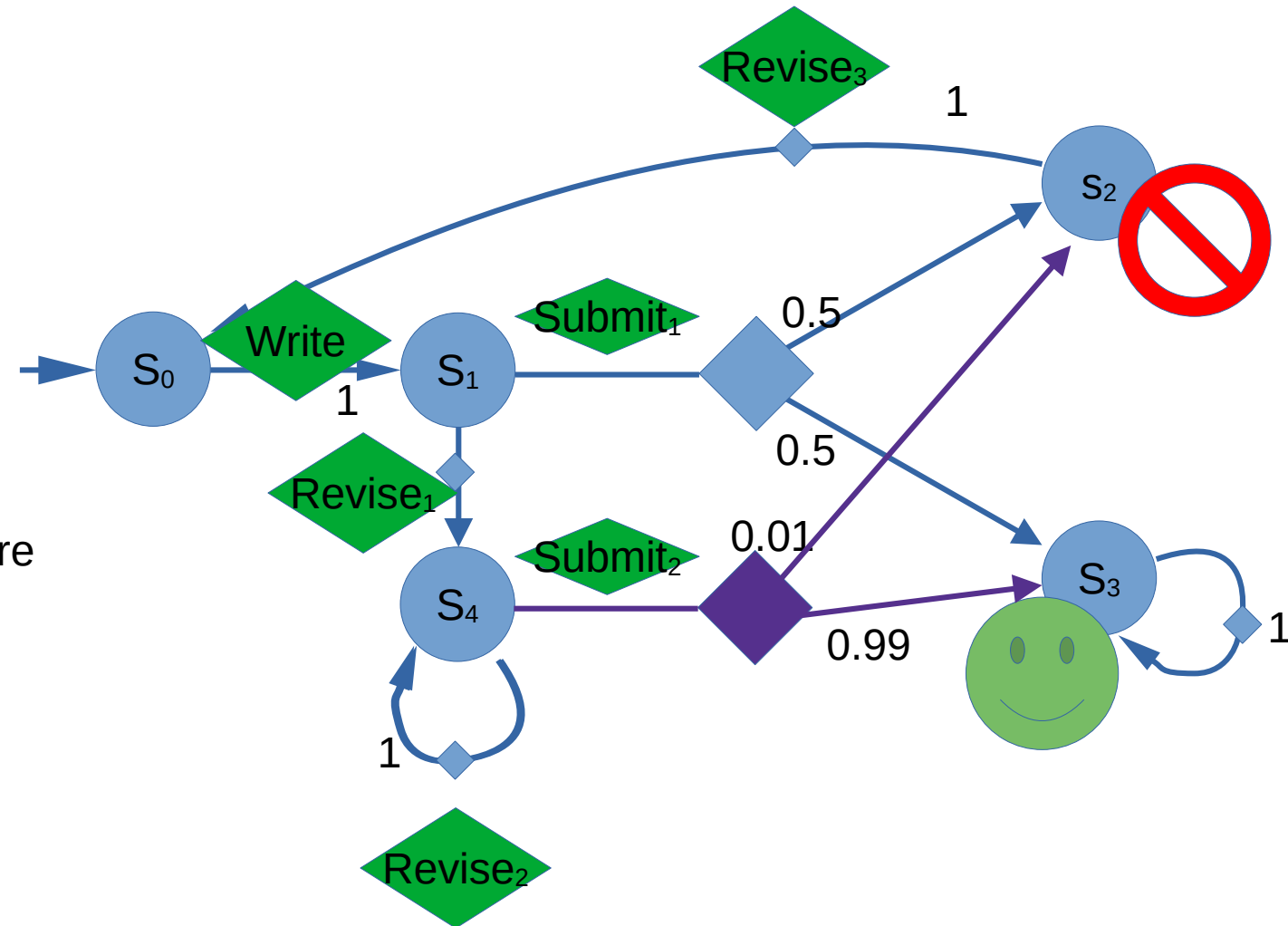
What is the optimal **Strategy** s.t. we publish a paper?

Strategy:

$\sigma : Q \rightarrow \text{Dist}\{\text{Write}, \text{Submit}_x, \text{Revise}_x\}$

Theorem:

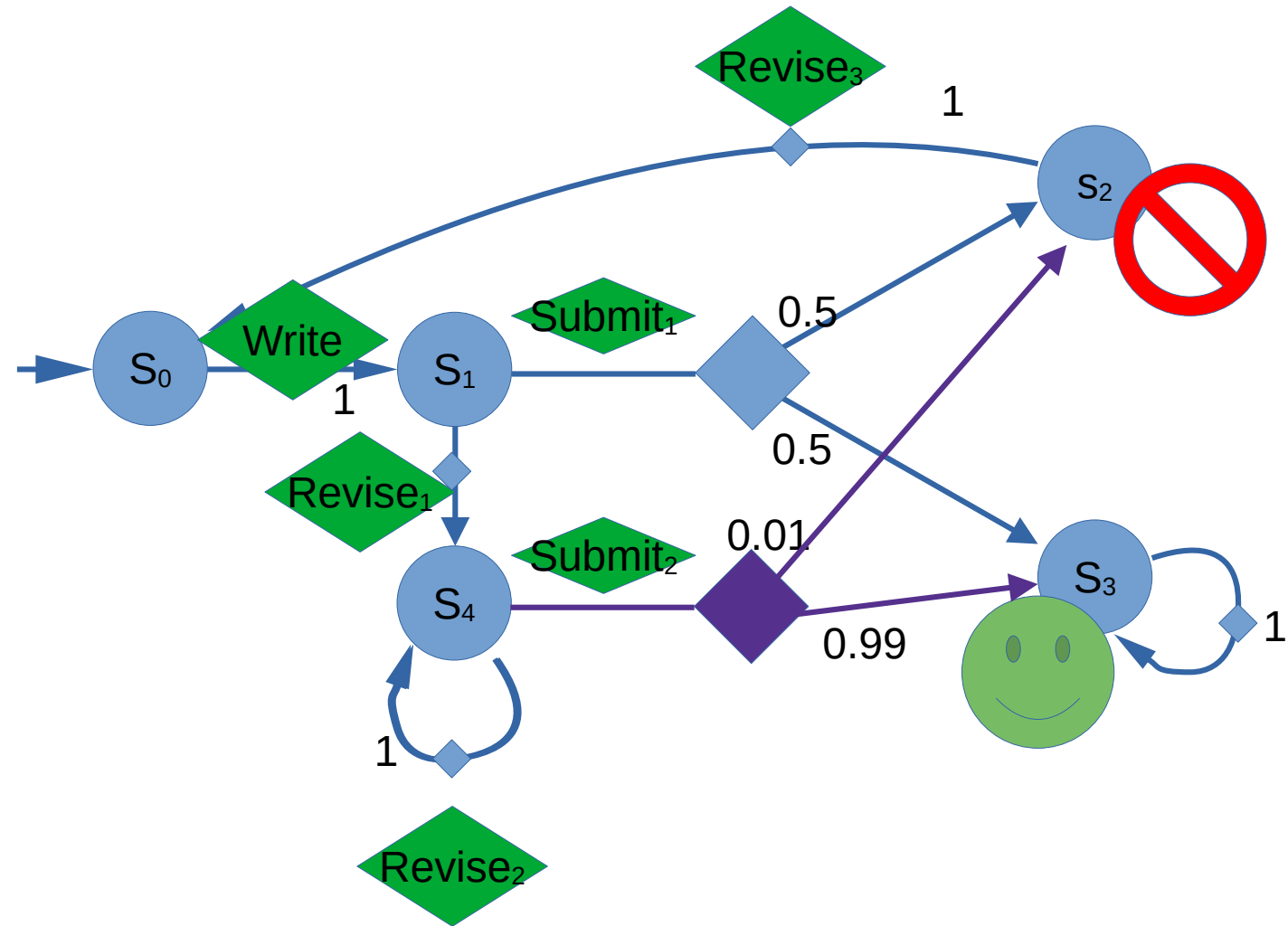
The optimal strategy can be computed and is defined by the Bellman equations. Furthermore a deterministic optimal strategy exists.



Q-learning

Model-free learning

- Transition probabilities are unknown
- Costs are unknown
- States are visible
- Actions are visible
- Works on samples of the system



Q-learning

Model-free learning

- Transition probabilities are unknown
- Costs are unknown
- States are visible
- Actions are visible
- Works on samples of the system

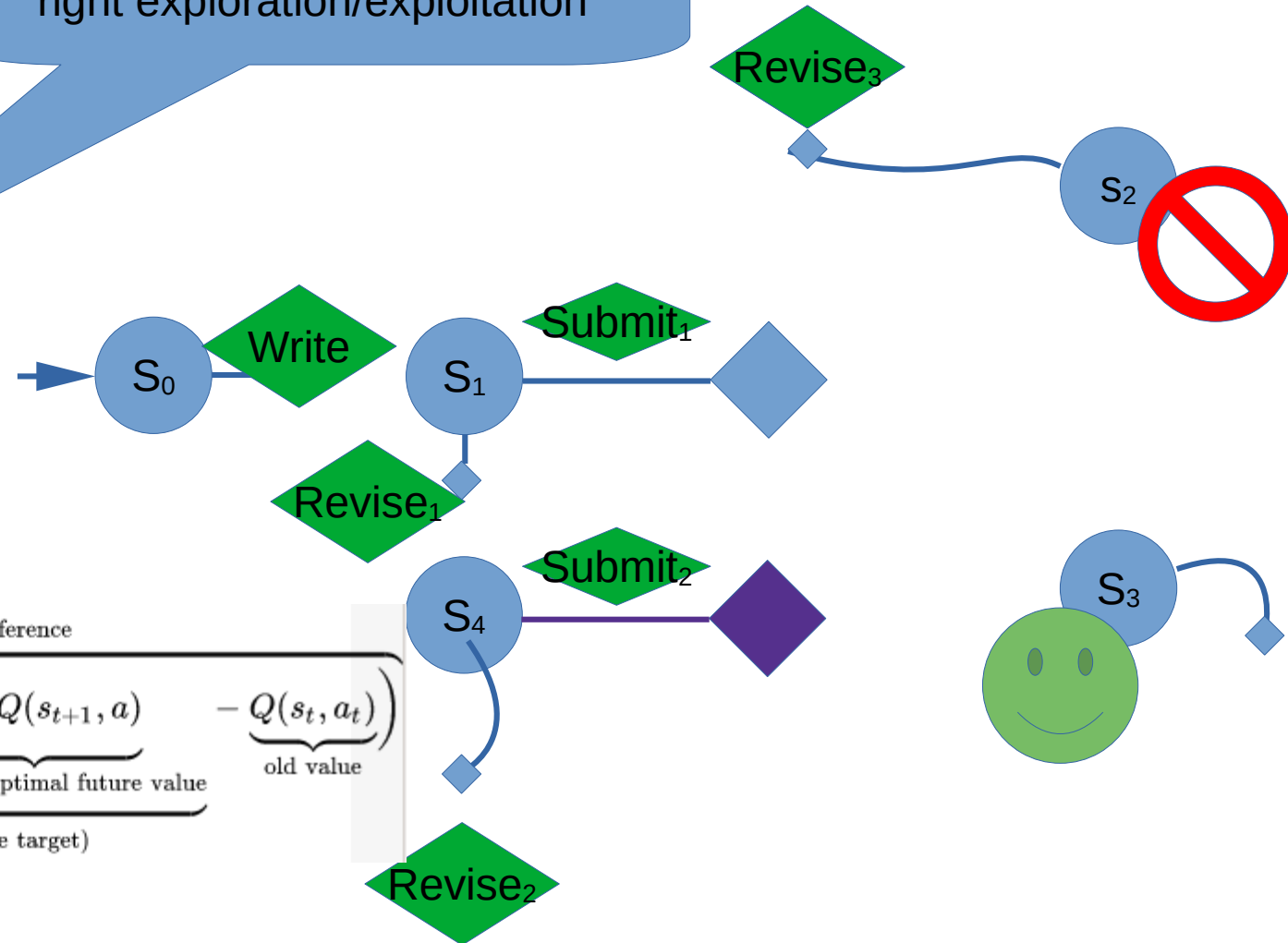
In practice, pick the 'right' α and γ + right exploration/exploitation

Provable convergence to discounted optimum

... if the learning-rate converges to zero in the limit AND all actions have non-zero probability of being sampled.

$$Q^{new}(s_t, a_t) \leftarrow \underbrace{Q(s_t, a_t)}_{\text{old value}} + \underbrace{\alpha}_{\text{learning rate}} \cdot \underbrace{\left(\underbrace{r_t}_{\text{reward}} + \underbrace{\gamma}_{\text{discount factor}} \cdot \underbrace{\max_a Q(s_{t+1}, a)}_{\text{estimate of optimal future value}} - \underbrace{Q(s_t, a_t)}_{\text{old value}} \right)}_{\text{new value (temporal difference target)}}$$

temporal difference



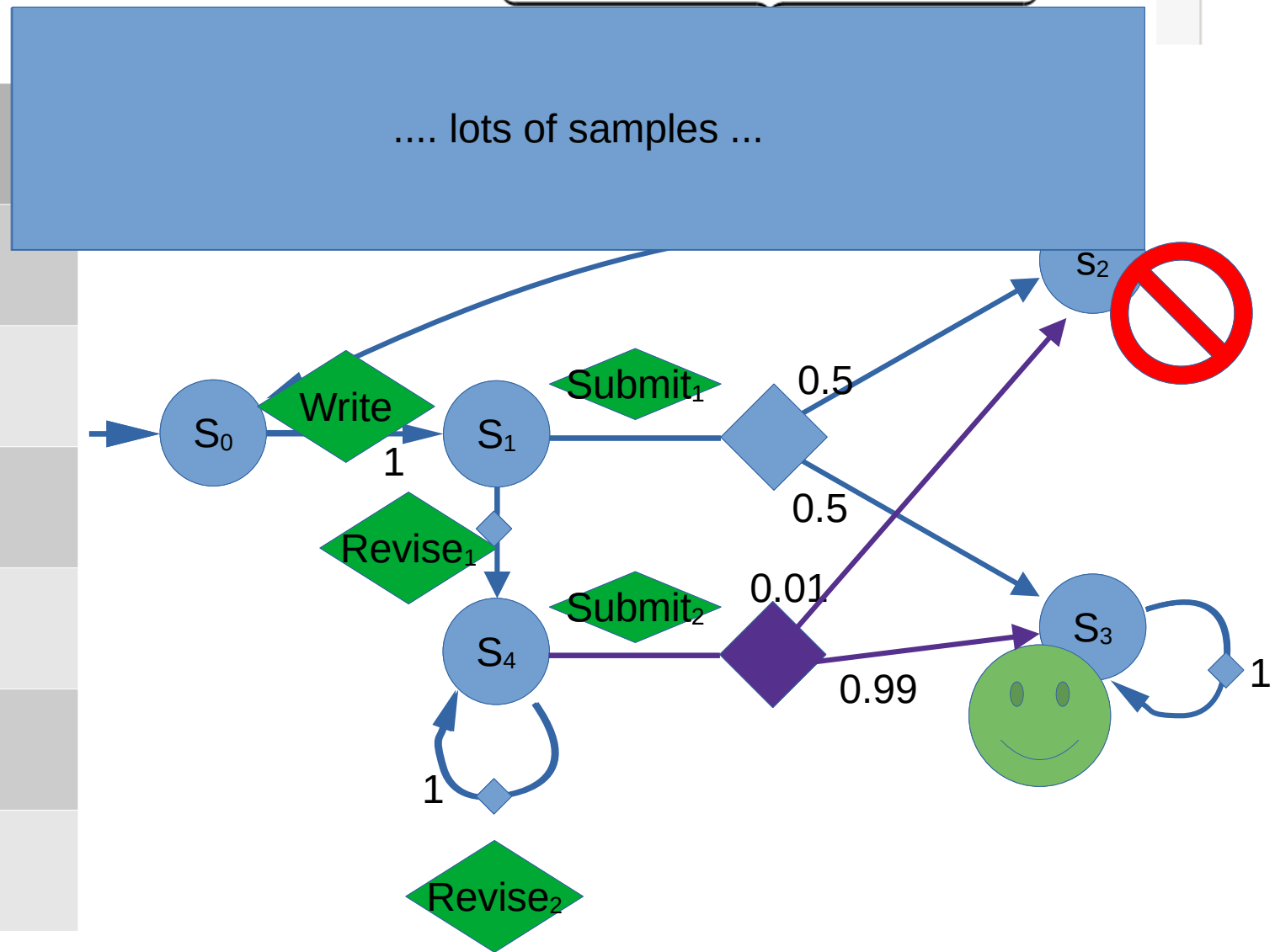
Assume action-reward -1, goal S_4 ,
Discount=0.9, learning-rate=0.5,
Greedy exploitation

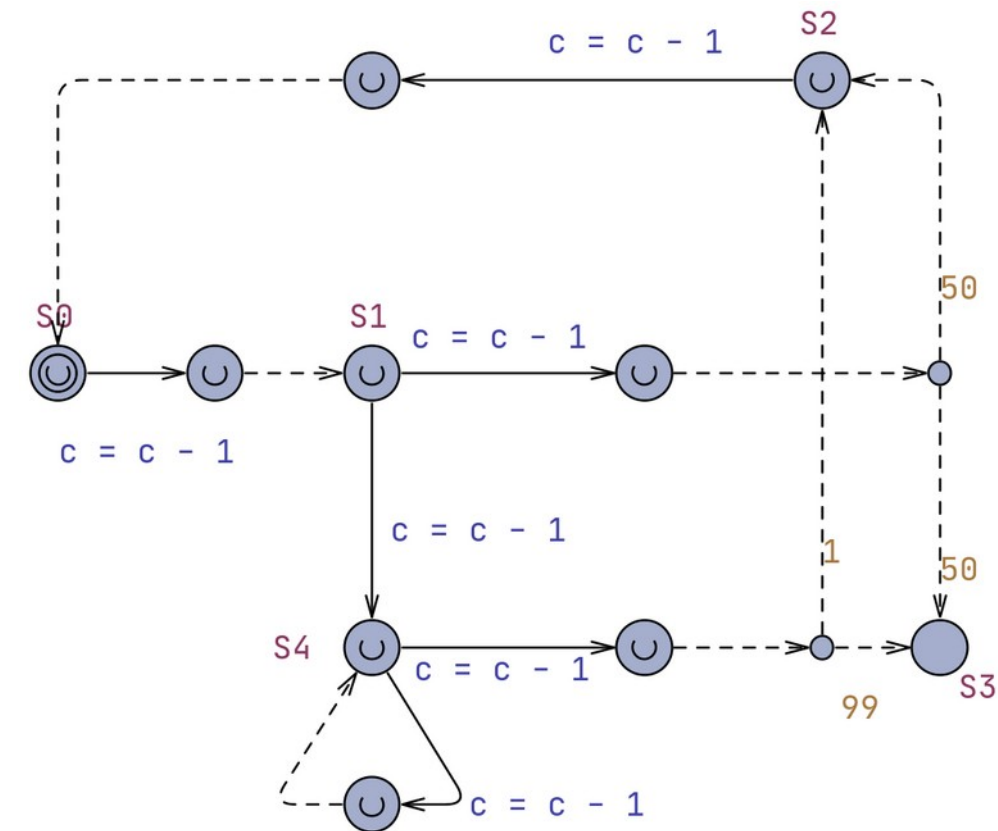
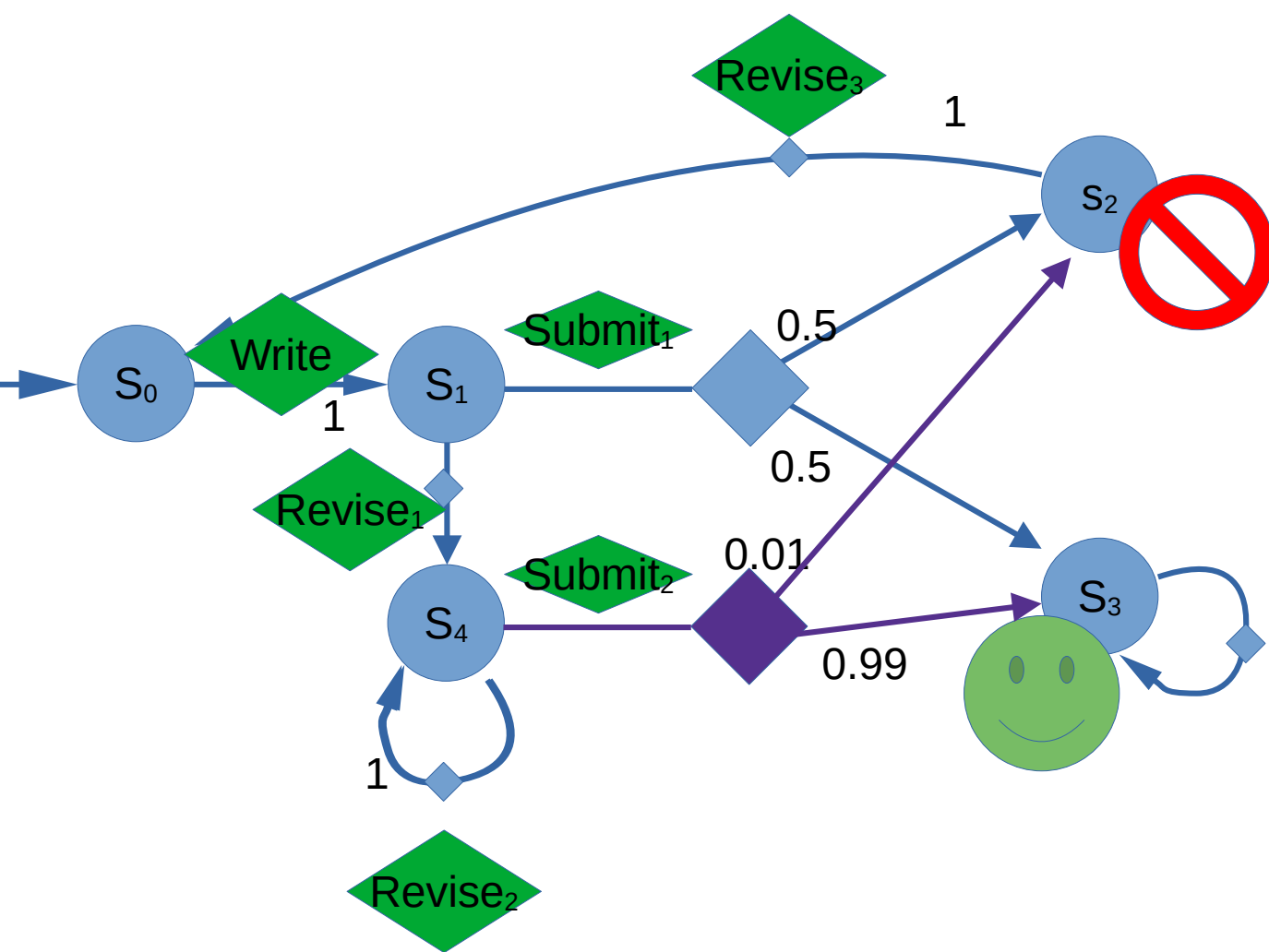
$$Q^{new}(s_t, a_t) \leftarrow \underbrace{Q(s_t, a_t)}_{\text{old value}} + \underbrace{\alpha}_{\text{learning rate}} \cdot \left(\underbrace{r_t}_{\text{reward}} + \underbrace{\gamma}_{\text{discount factor}} \cdot \underbrace{\max_a Q(s_{t+1}, a)}_{\text{estimate of optimal future value}} - \underbrace{Q(s_t, a_t)}_{\text{old value}} \right)$$

temporal difference

$Q(S_n,$ action) =

S_0	Write	-0.75
S_1	Submit ₁	-0.5
S_1	Revise ₁	-0.75
S_2	Revise ₃	0
S_3		0
S_4	Submit ₂	-0.5
S_4	Revise ₂	-0.75





```

strategy s = maxE(c) [<=100]: <> S3 // compute the strategy
saveStrategy(s, "filepath.json") // save the strategy
E[<=100;100] (max:c) under s // evaluate strategy performance
E[<=100;100] (max:c) // performance of random controller

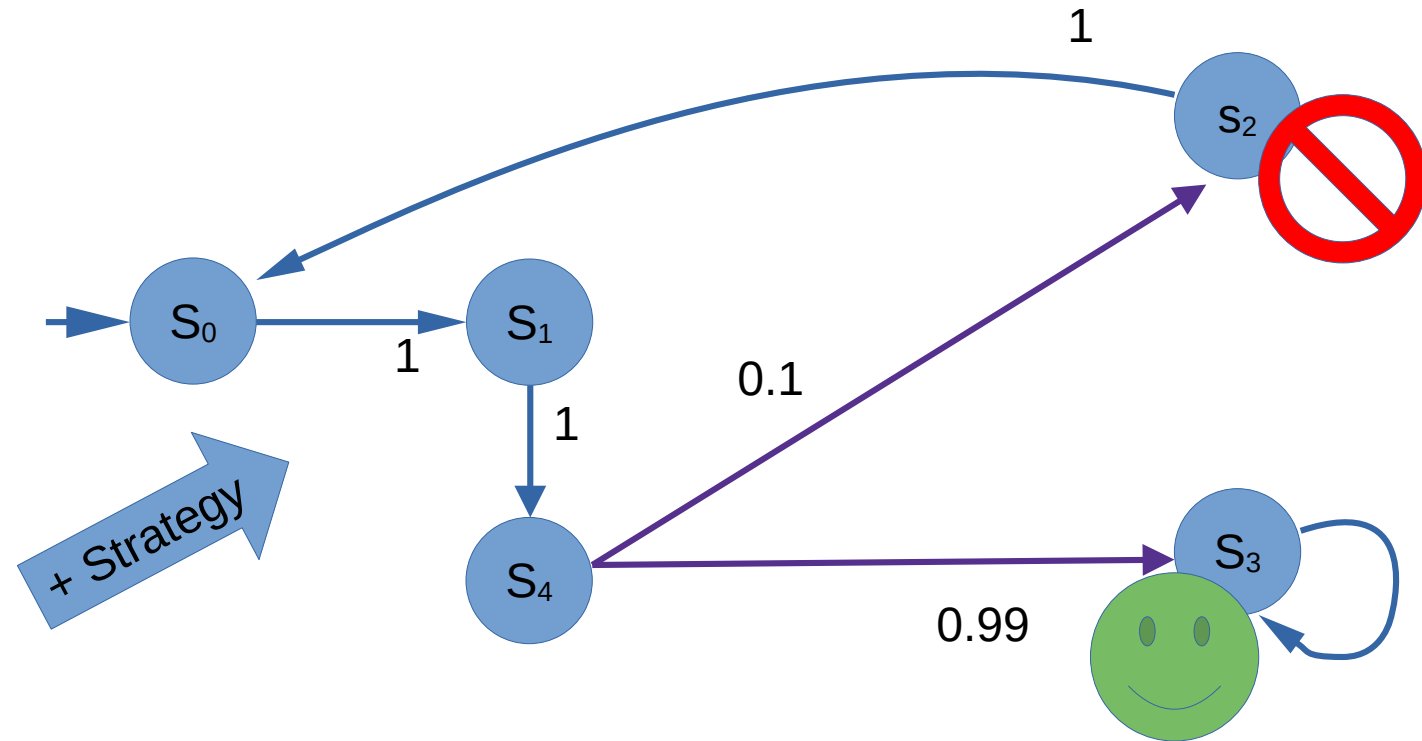
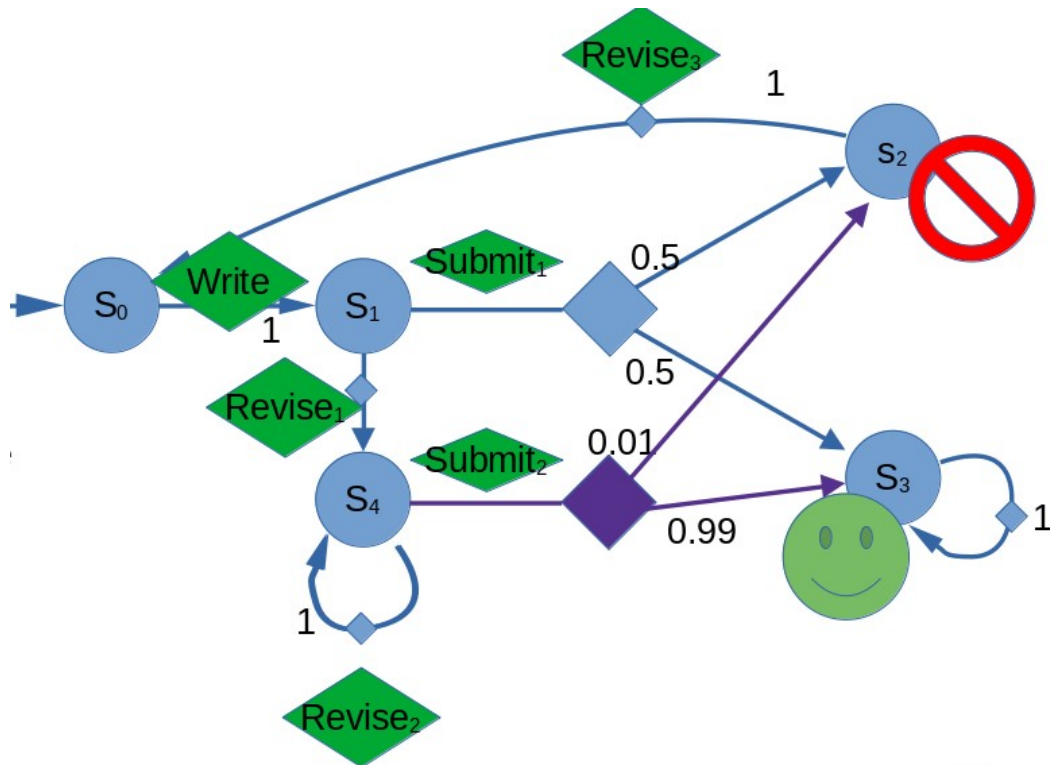
```


Q-Learning

Q-table gives us a strategy

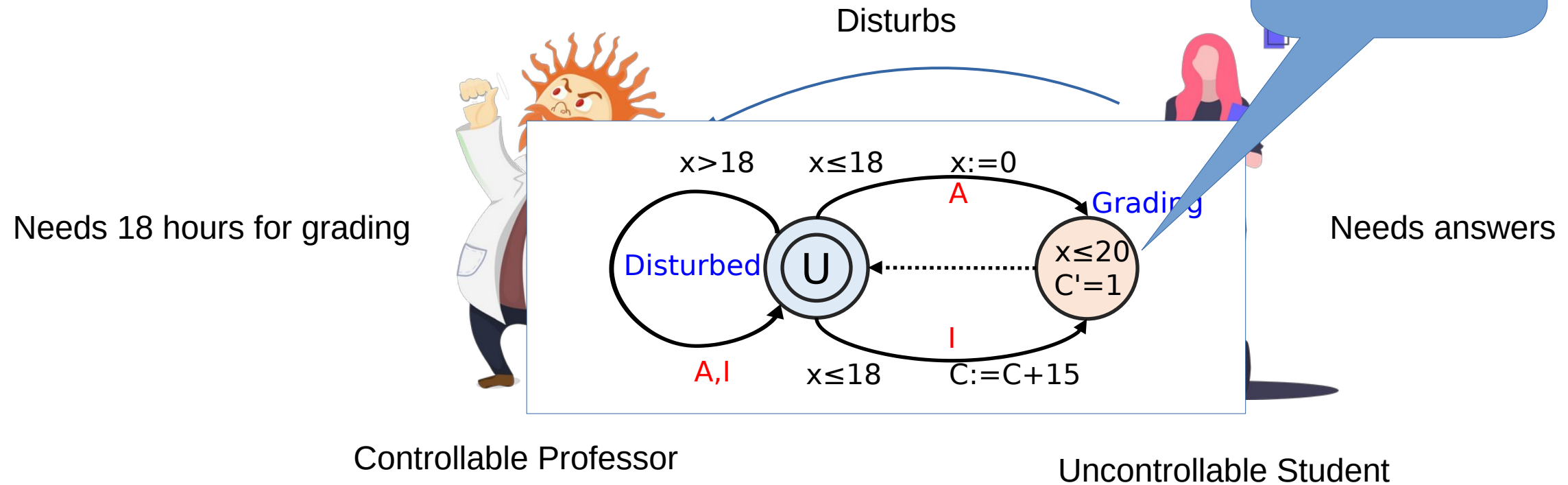
- For each state, pick "best" action

Strategy(Q-table) + MDP $\Rightarrow$ MC



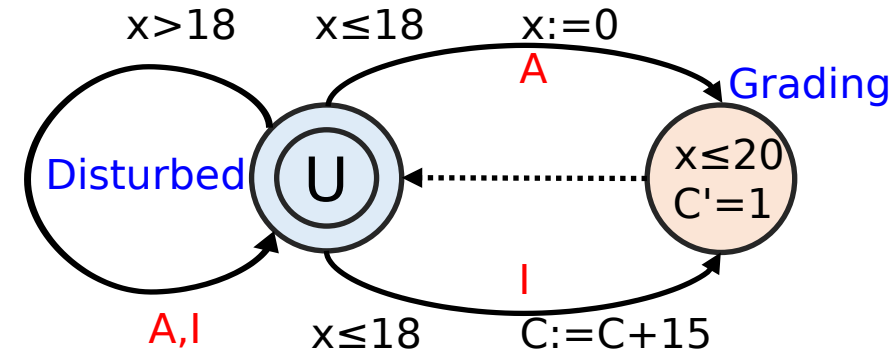
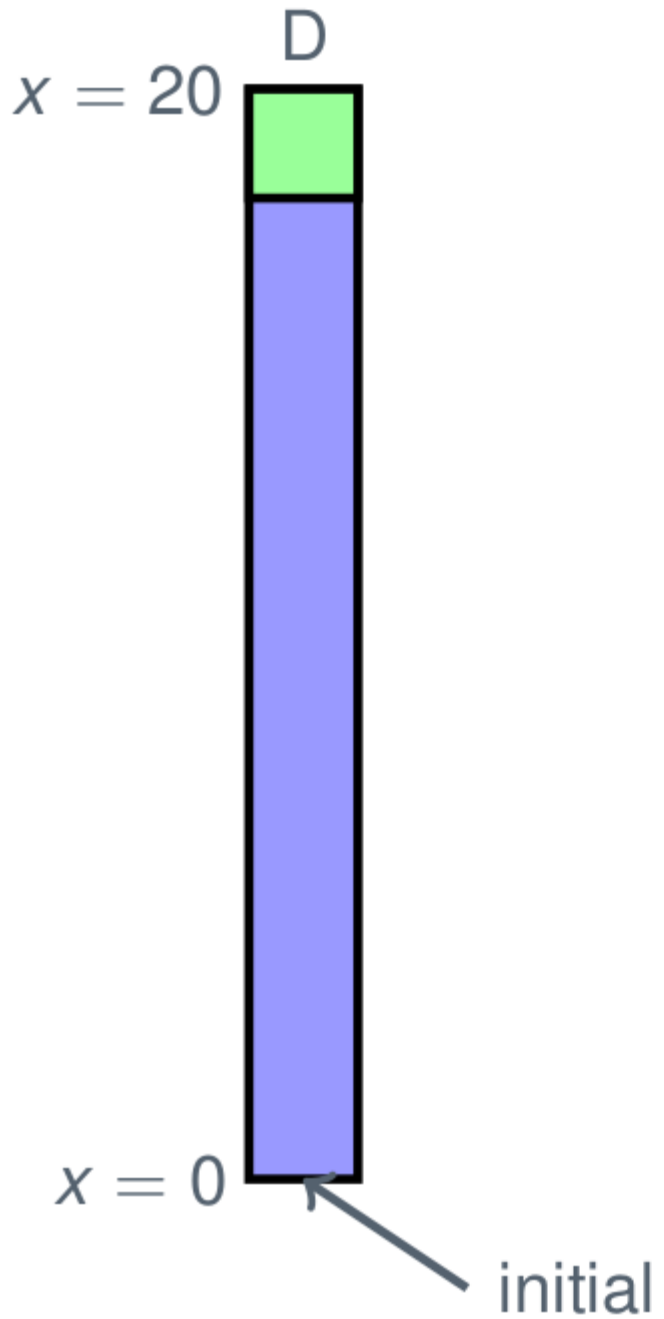
Euclidean Markov Decision Process

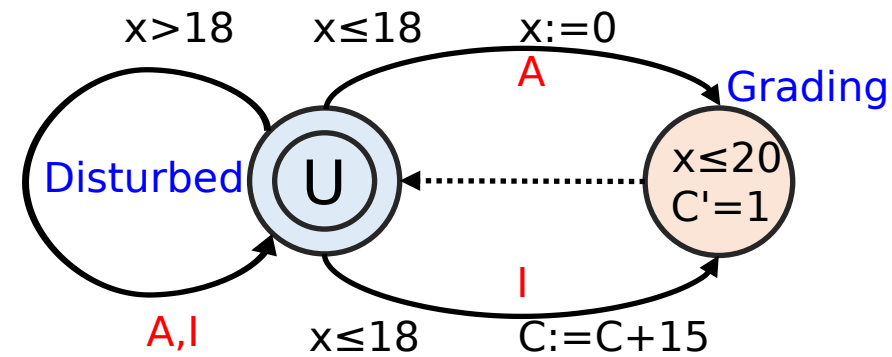
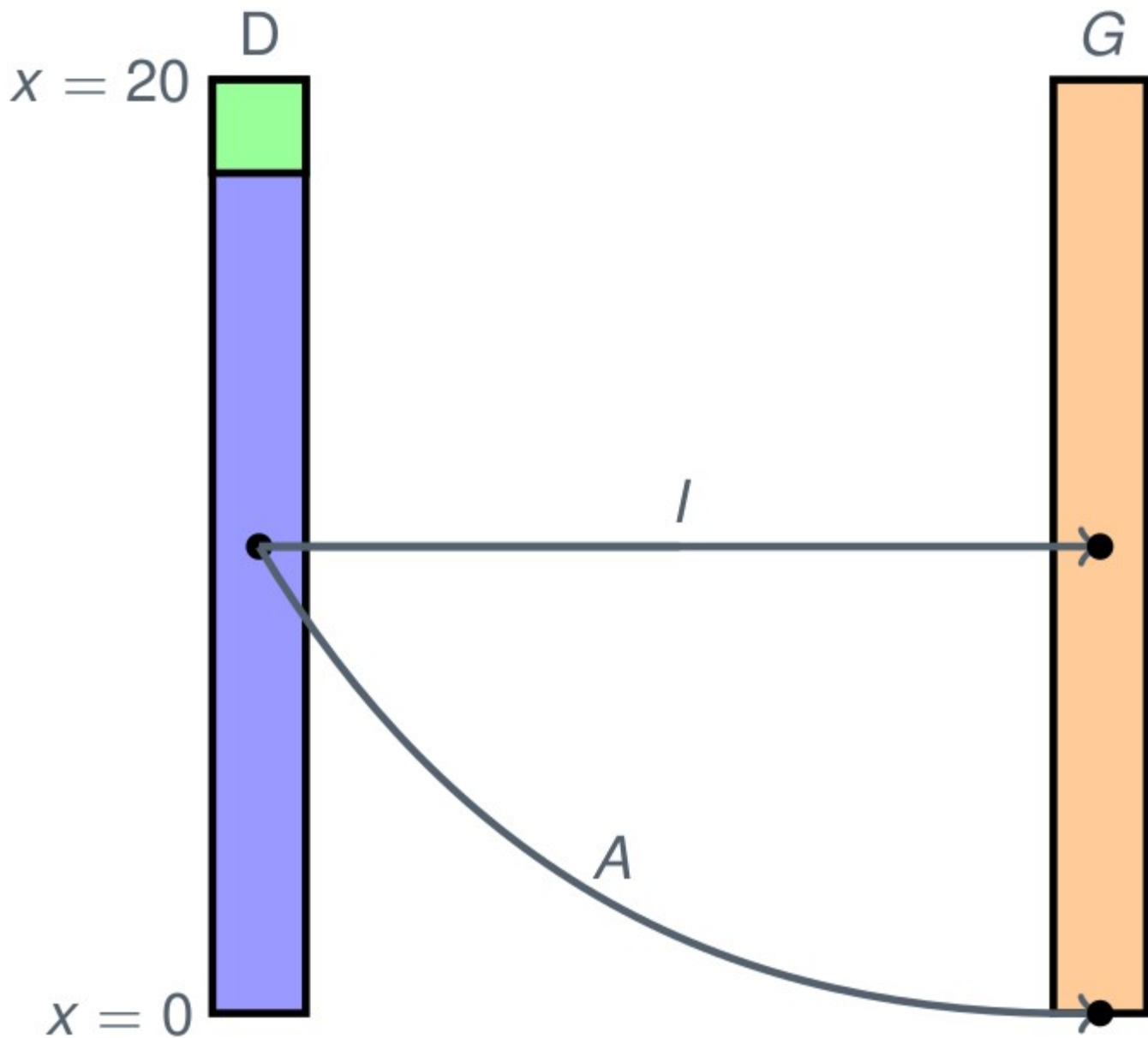
MDP With Continuous State Space

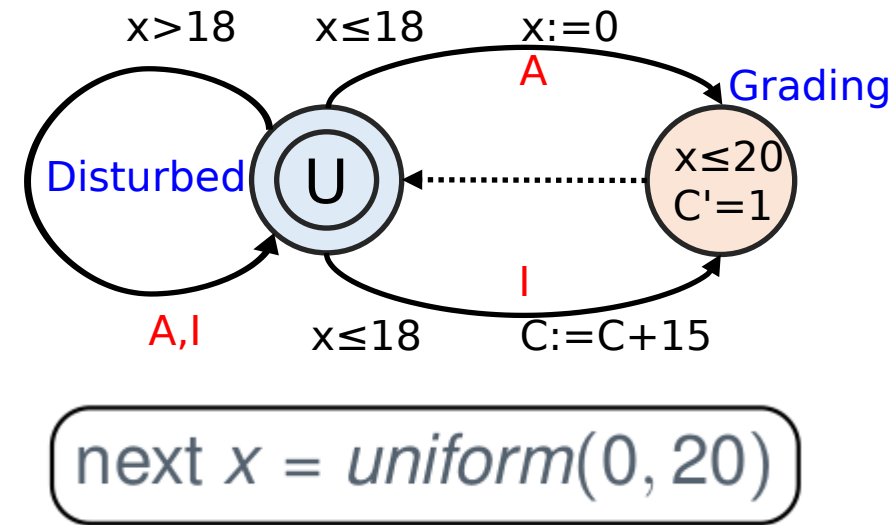
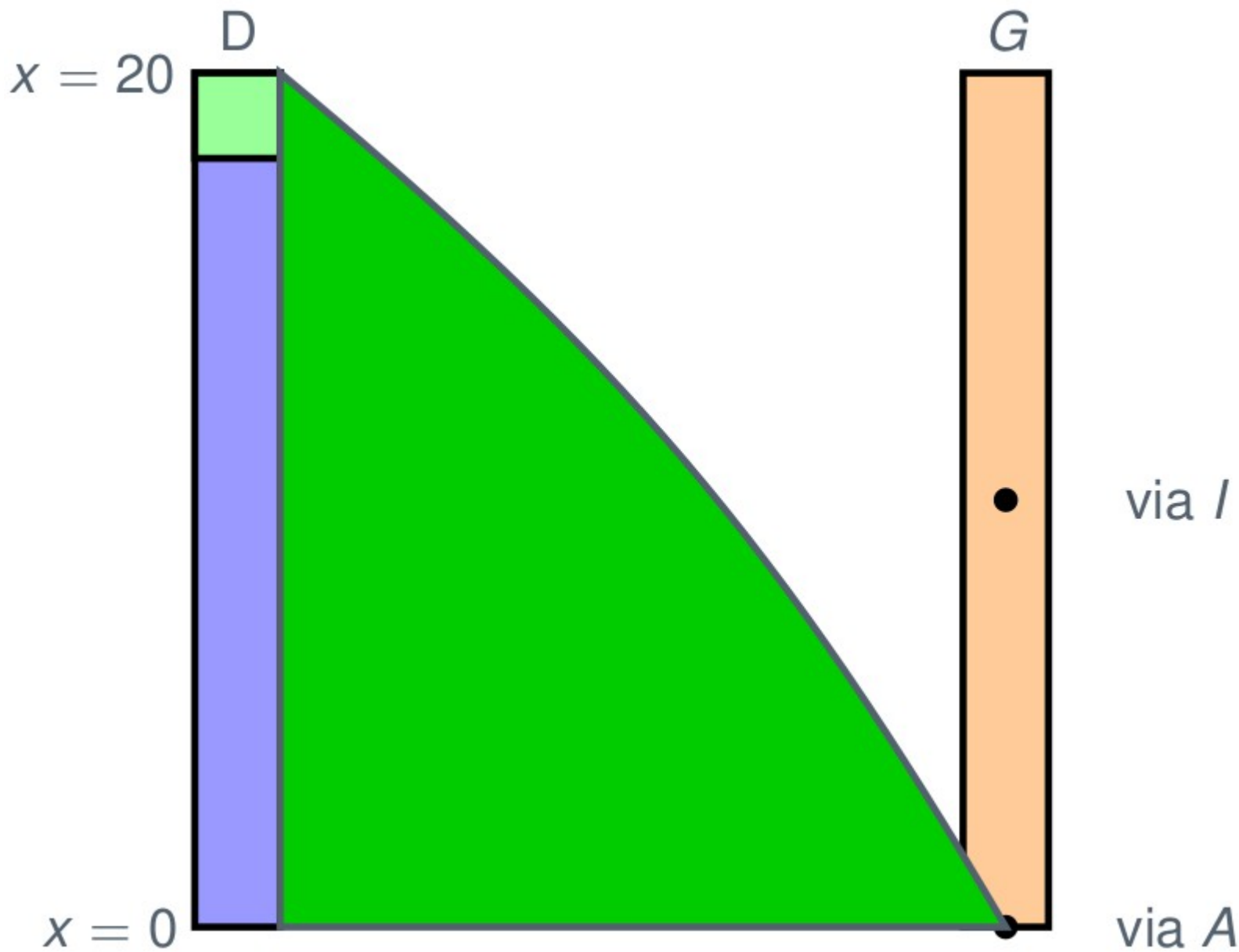


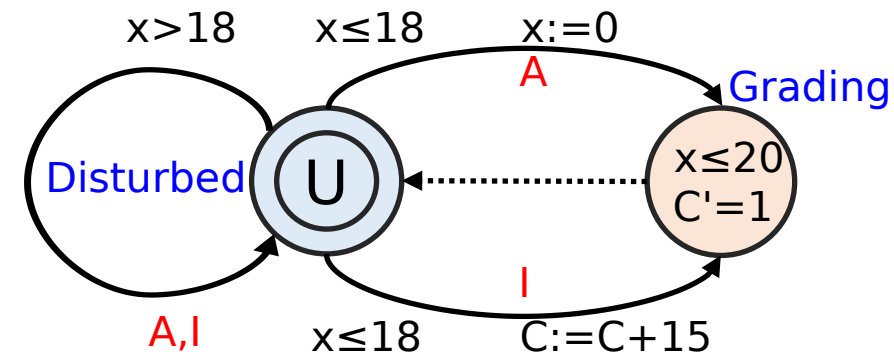
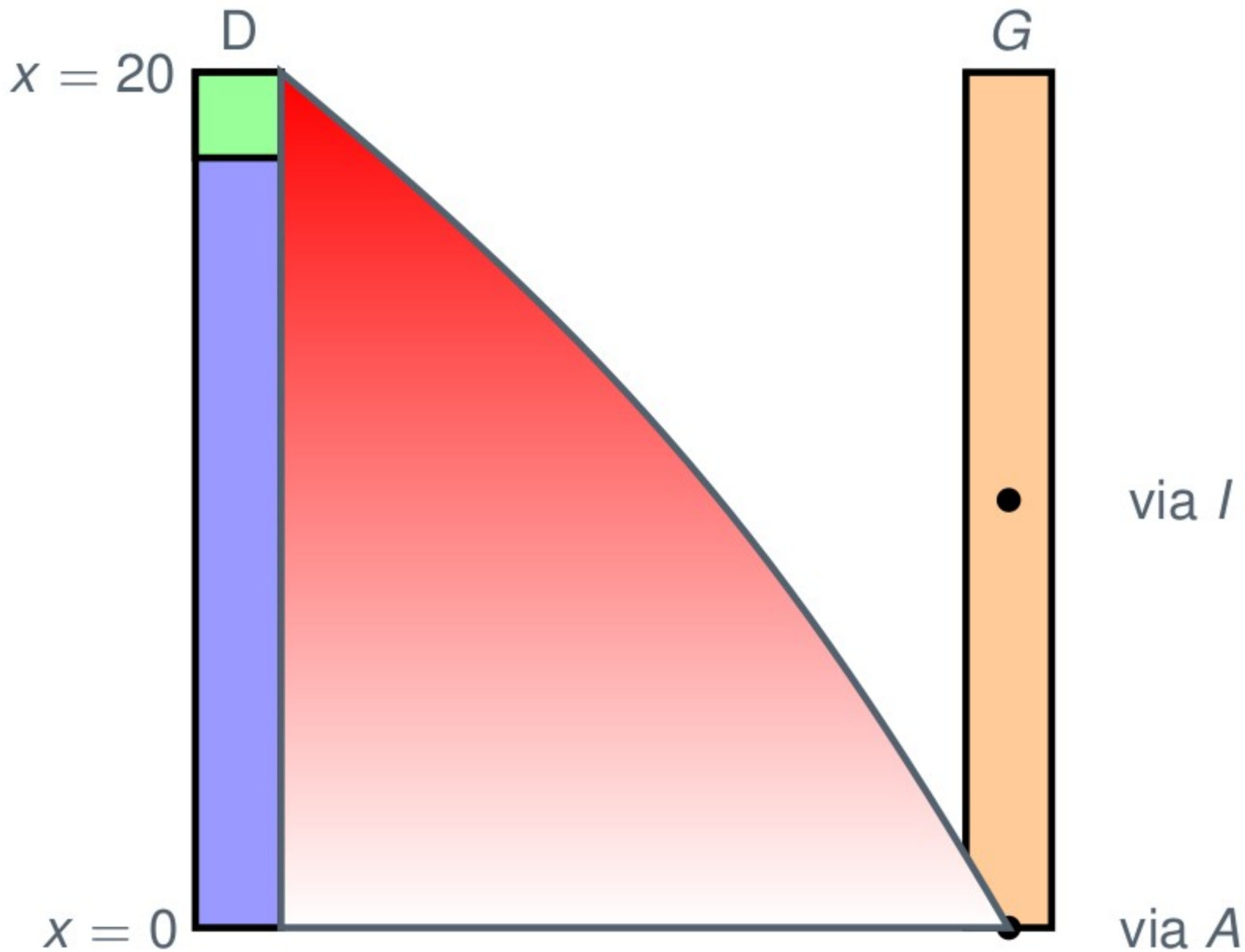
Choices:

- 1) Answer student, restart grading
- 2) Ignore the student, spend 15 hours on complaint later



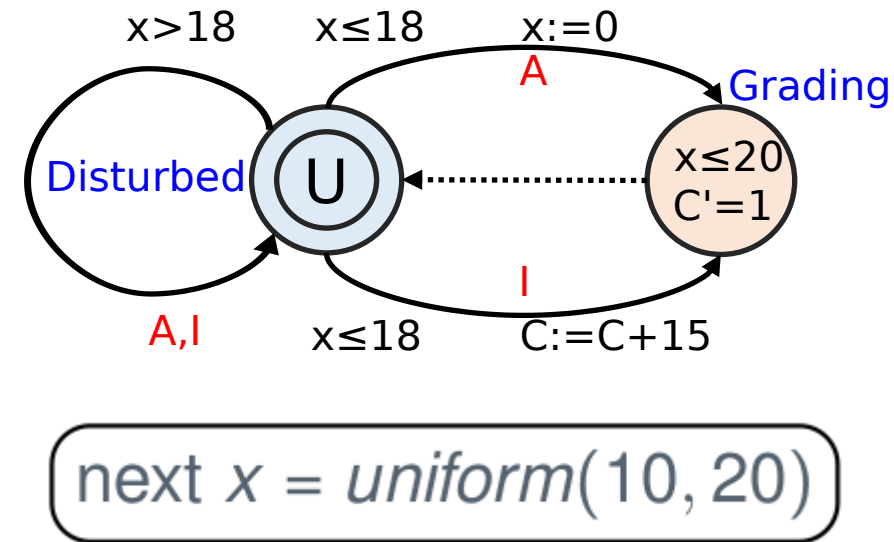
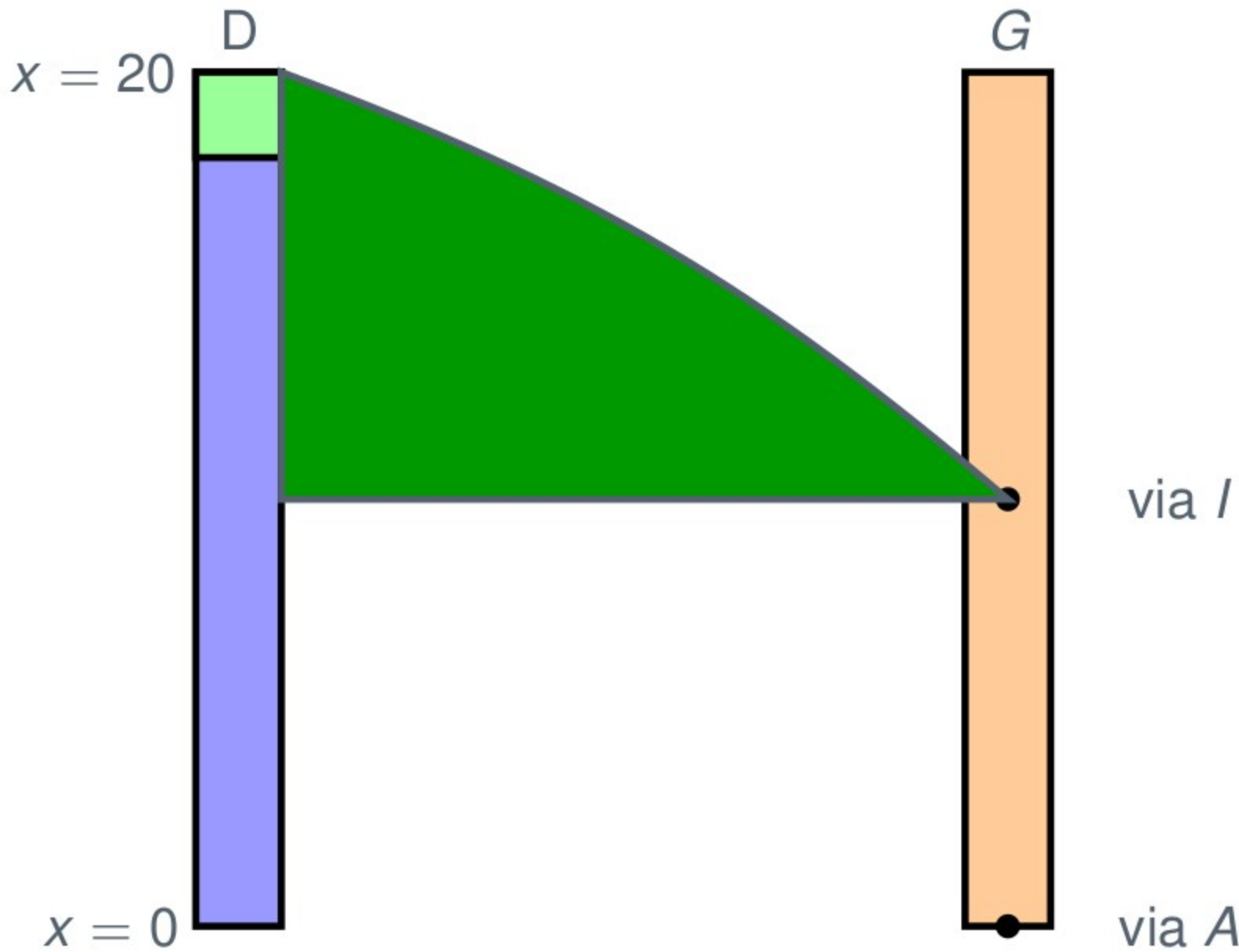


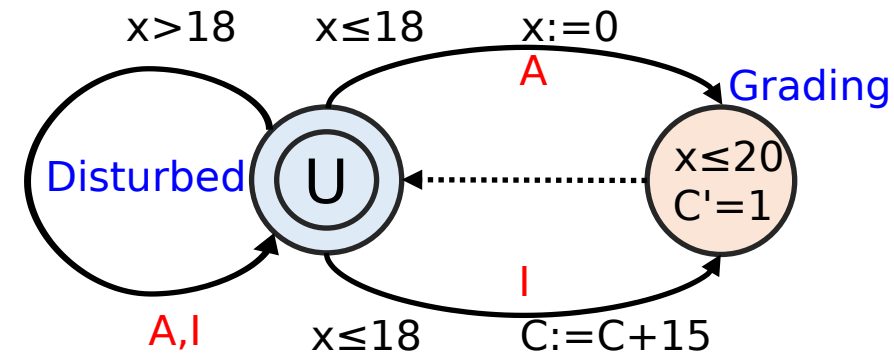
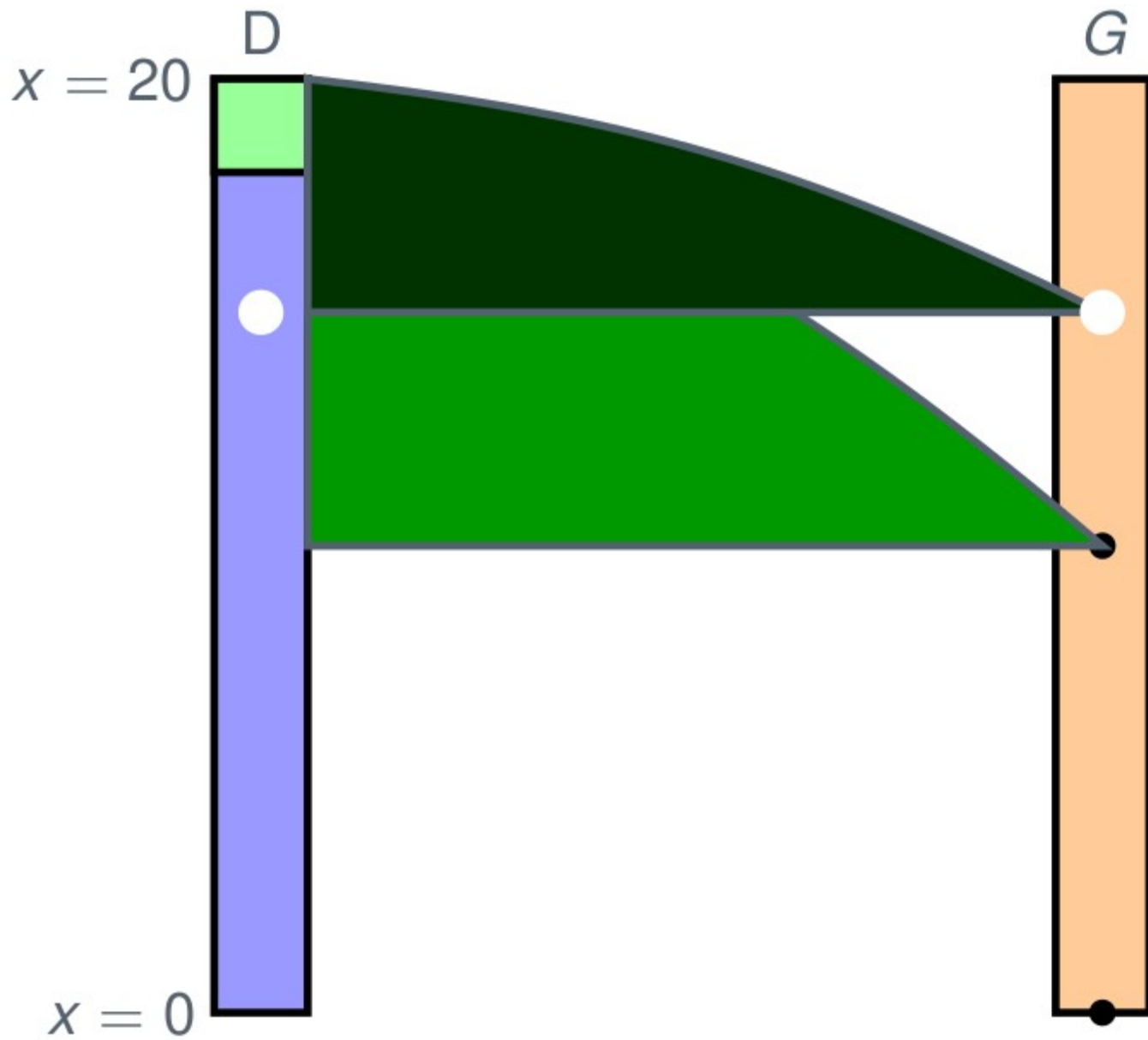




next $x = \text{uniform}(0, 20)$

next $C = C + \text{next } x$



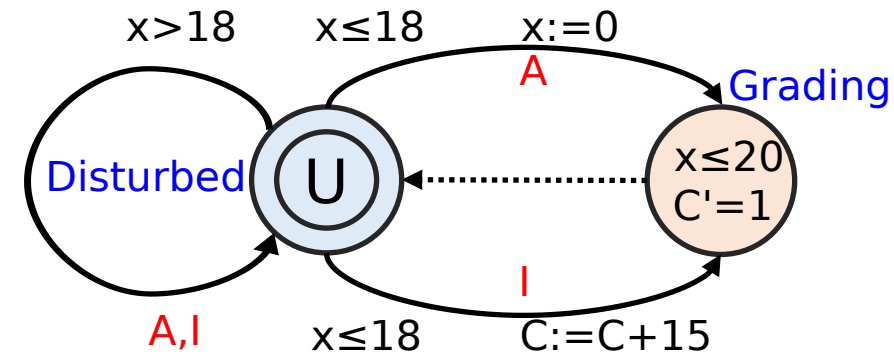
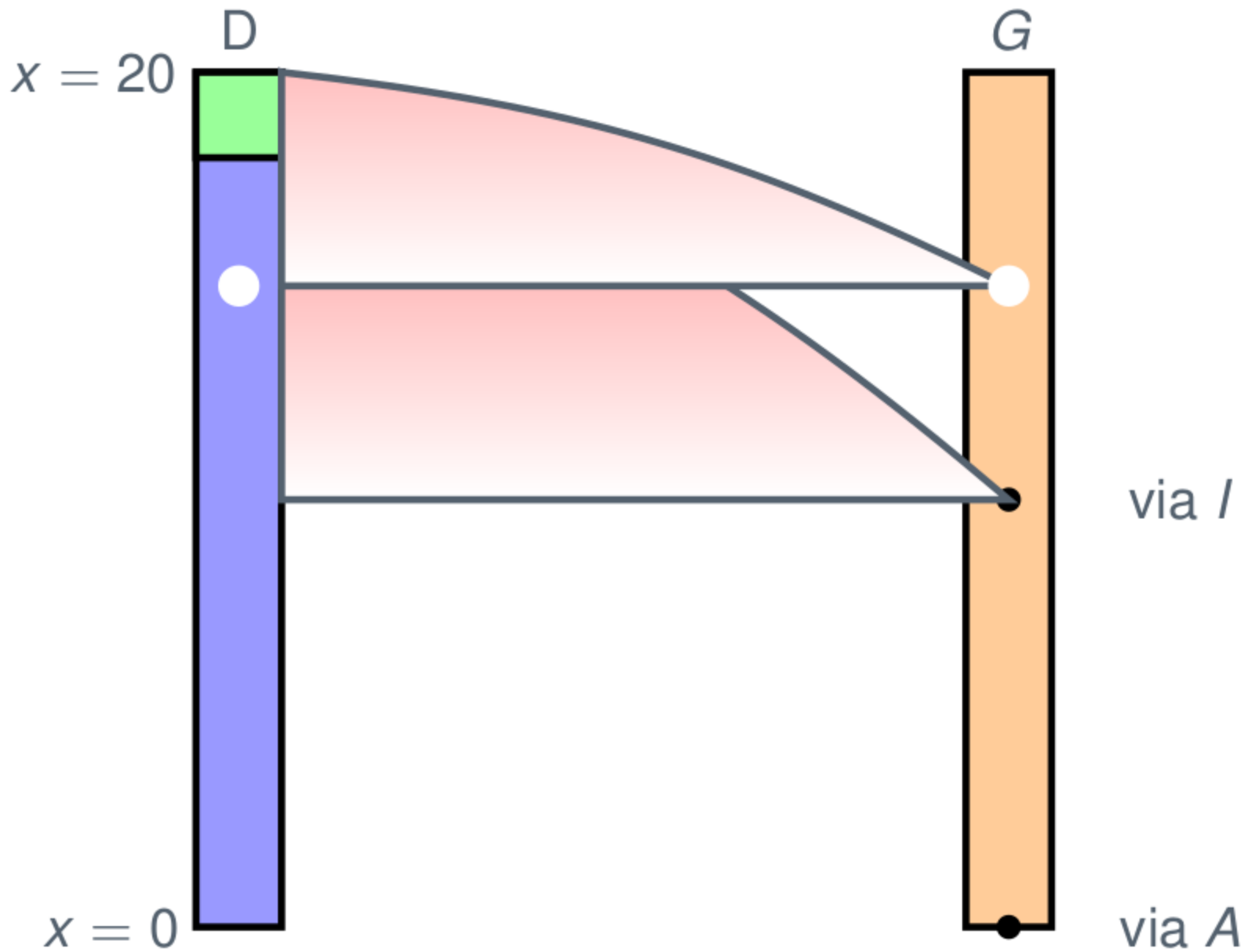


next $x = \text{uniform}(10, 20)$

next $x = \text{uniform}(15, 20)$

via I

via A



next $C = C + \text{next } x$

relative to x

Euclidean Markov Decision Process

MDP With Continuous State Space

Undecidable to solve in theory and practice

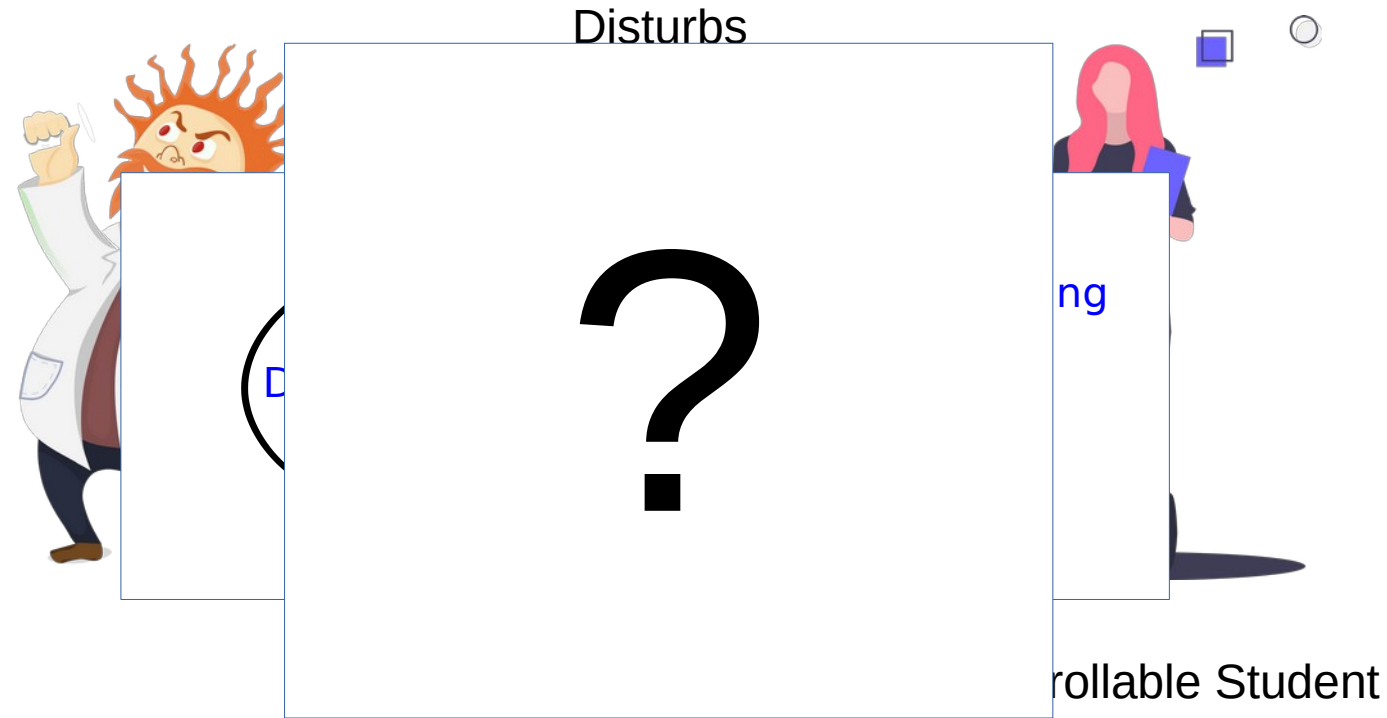
- Infinite (uncountable) number of states
- We cannot map to a regular MDP

Unobservable

- No direct cost function
- No direct transition function

Solution

- Discretize!
- What is a good general discretization?



$$Q(\nu) = \min(Q(\nu, A), Q(\nu, I))$$



$$Q(\nu, A)$$



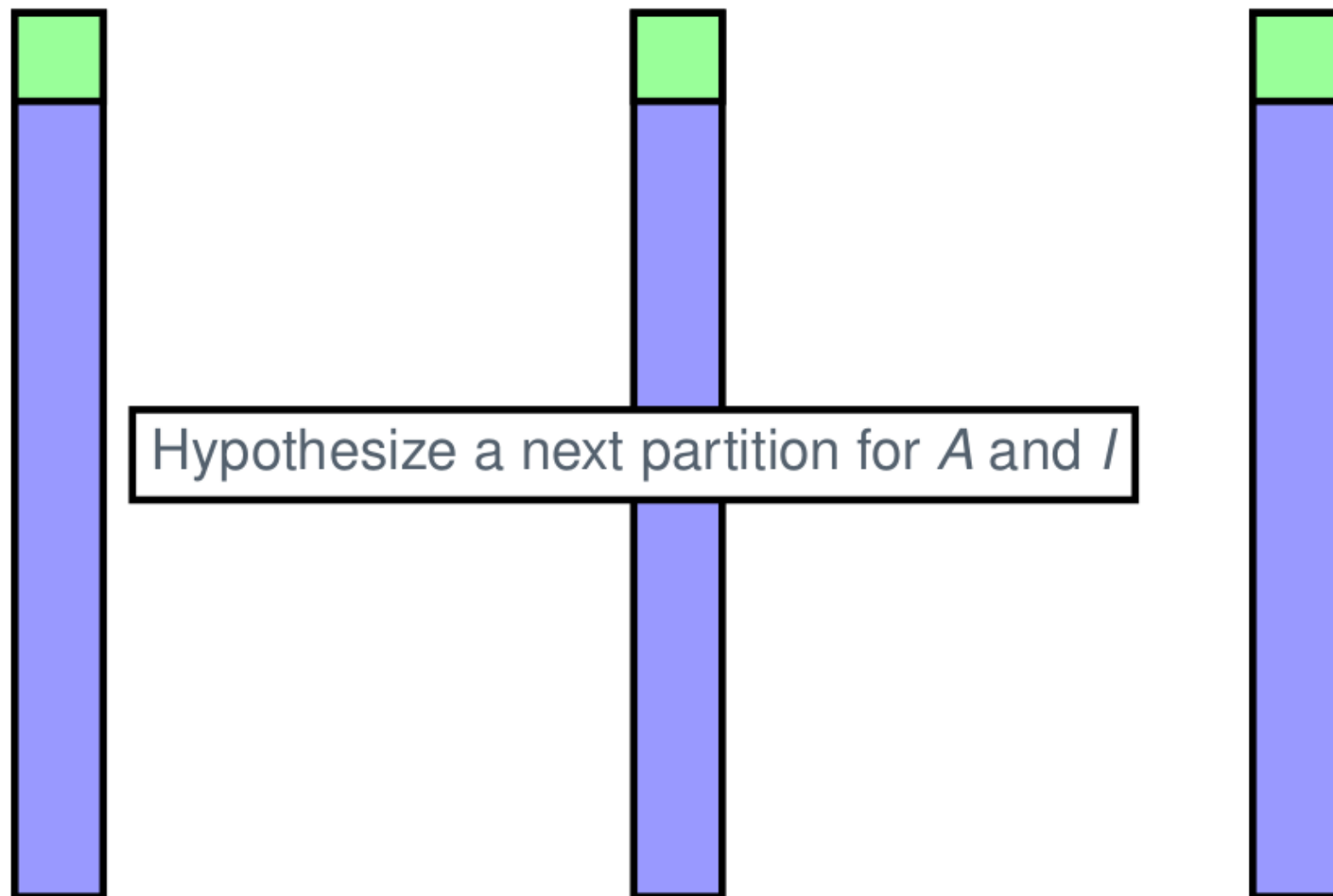
$$Q(\nu, I)$$



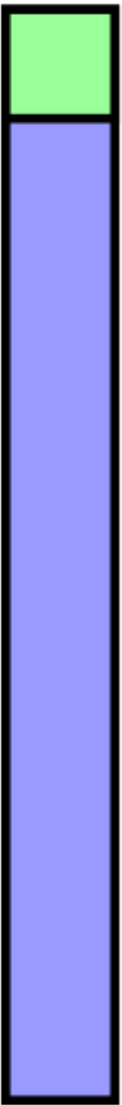
$$Q(\nu) = \min(Q(\nu, A), Q(\nu, I))$$

$$Q(\nu, A)$$

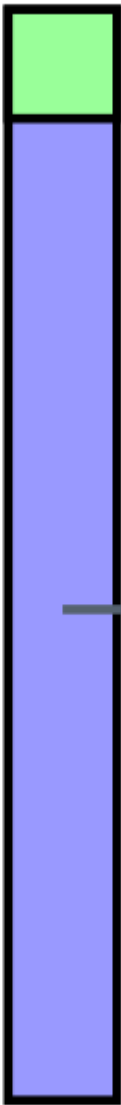
$$Q(\nu, I)$$



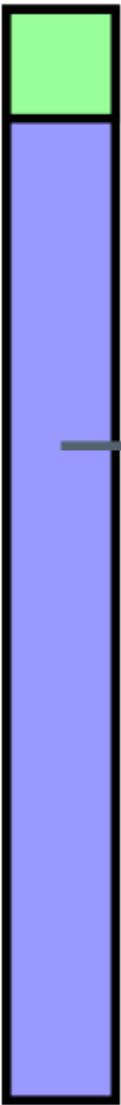
$Q(\nu) = \min(Q(\nu, A), Q(\nu, I))$



$Q(\nu, A)$



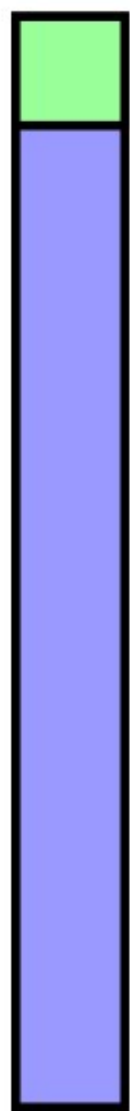
$Q(\nu, I)$



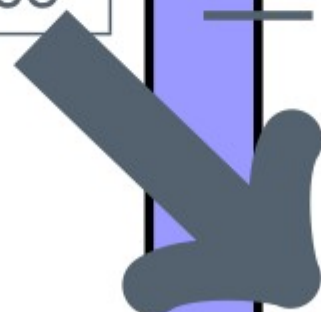
$$Q(\nu) = \min(Q(\nu, A), Q(\nu, I))$$

$$Q(\nu, A)$$

$$Q(\nu, I)$$



σ'	Mean
f'	Frequency
v'	Variance



$$\begin{array}{c|c} \sigma' & \\ f' & 0 \\ v' & 0 \end{array}$$

$$x = 9$$

$$\begin{array}{c|c} \sigma' & \\ f' & 0 \\ v' & 0 \end{array}$$



$$\begin{array}{c|c} \sigma' & \\ f' & 0 \\ v' & 0 \end{array}$$

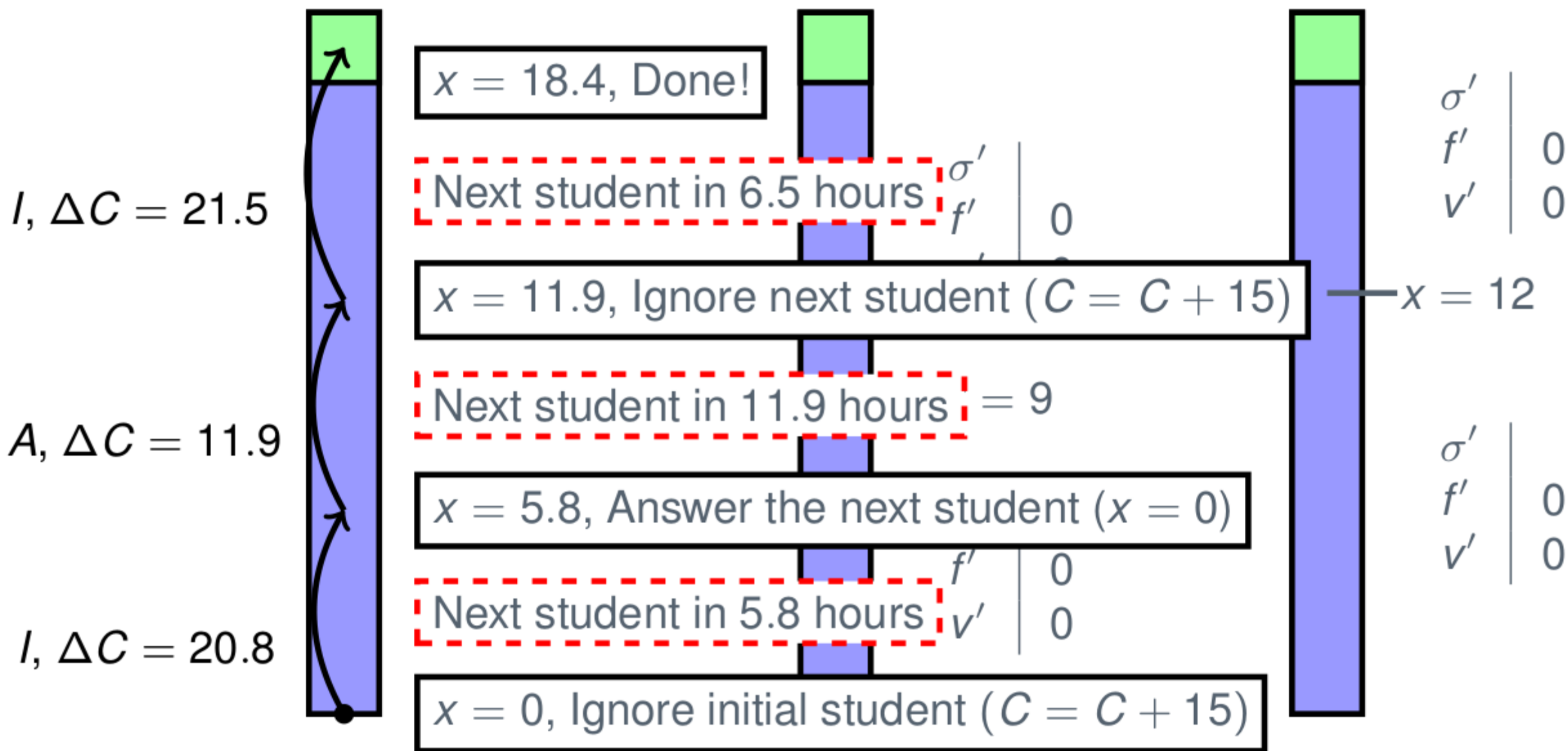
$$x = 12$$

$$\begin{array}{c|c} \sigma' & \\ f' & 0 \\ v' & 0 \end{array}$$

$$Q(\nu) = \min(Q(\nu, A), Q(\nu, I))$$

$$Q(\nu, A)$$

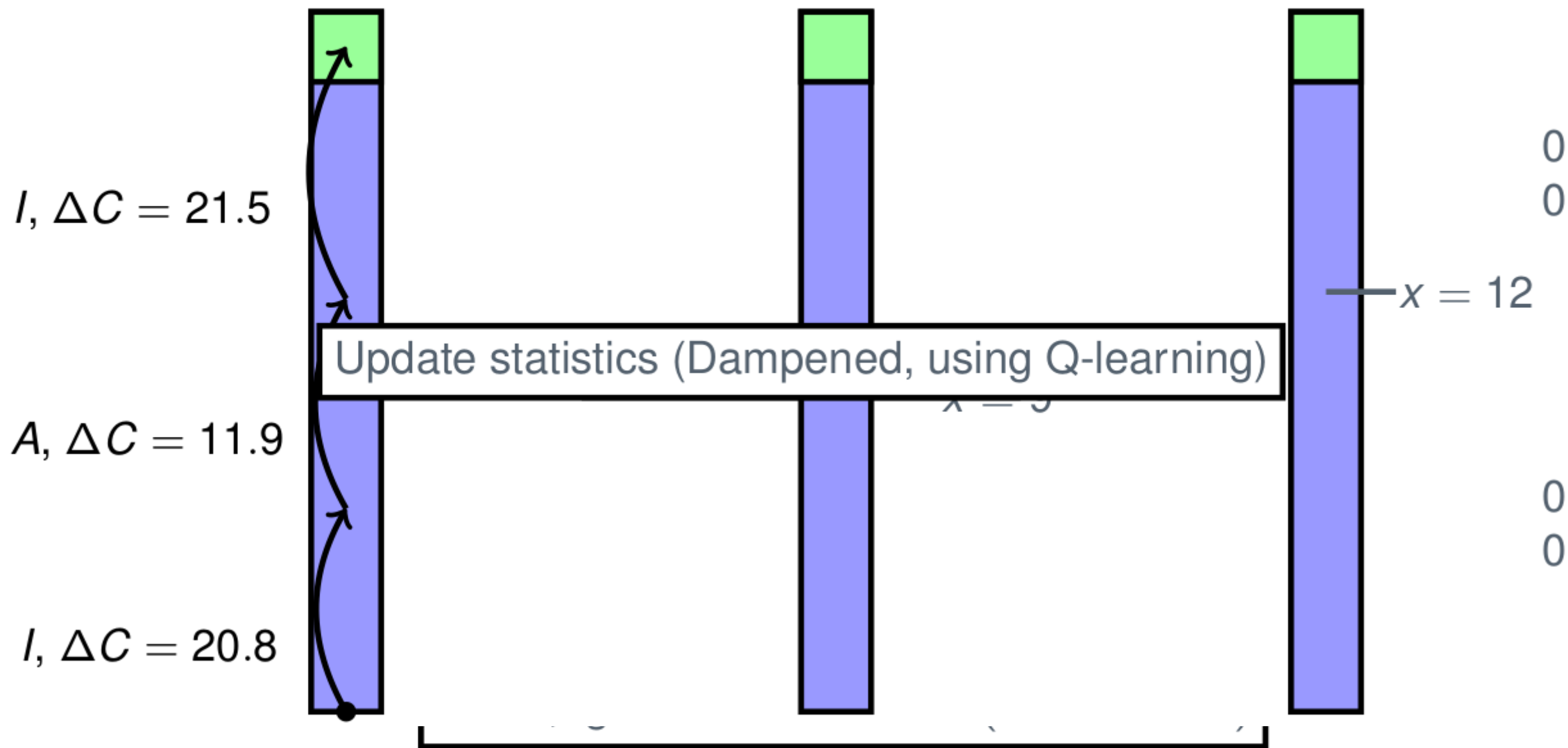
$$Q(\nu, I)$$



$$Q(\nu) = \min(Q(\nu, A), Q(\nu, I))$$

$$Q(\nu, A)$$

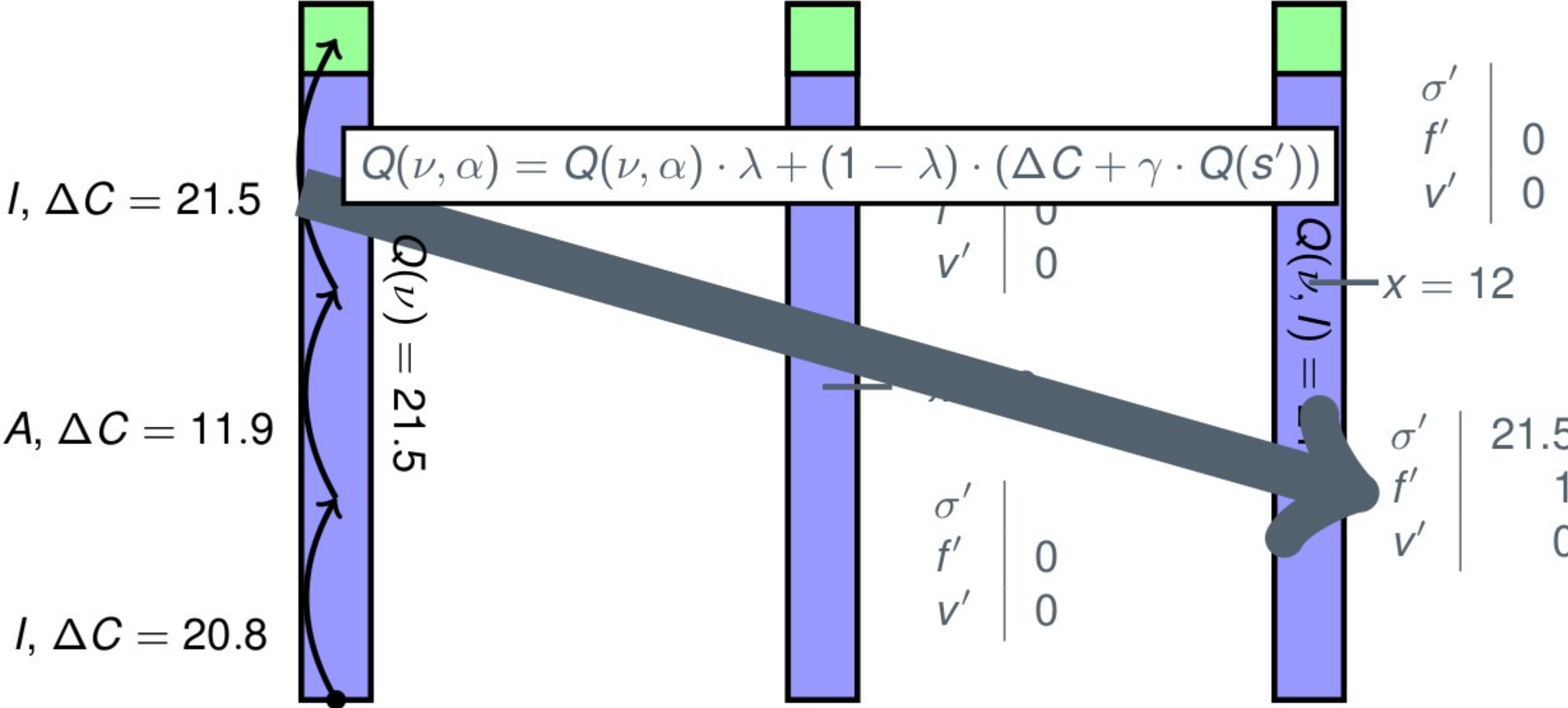
$$Q(\nu, I)$$



$$Q(\nu) = \min(Q(\nu, A), Q(\nu, I))$$

$$Q(\nu, A)$$

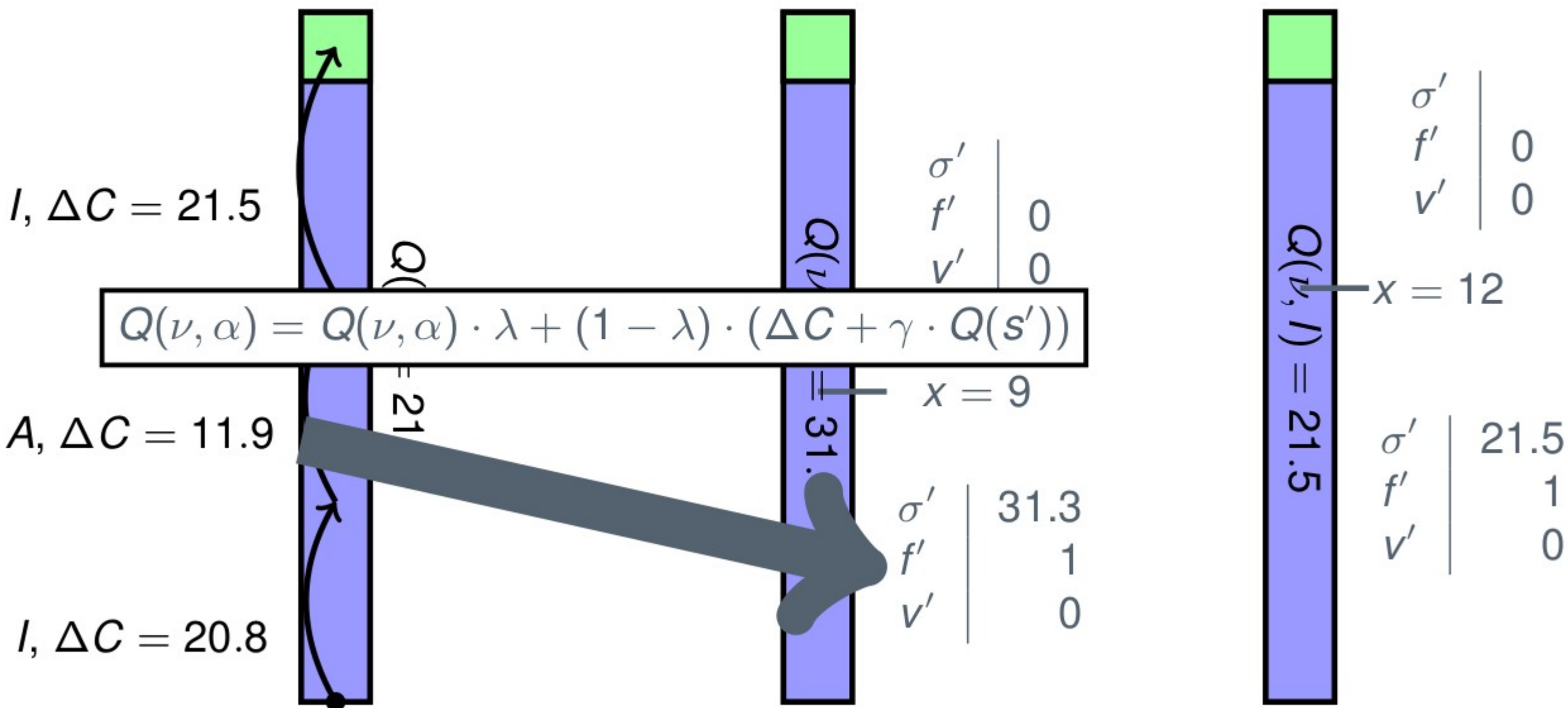
$$Q(\nu, I)$$



$$Q(\nu) = \min(Q(\nu, A), Q(\nu, I))$$

$$Q(\nu, A)$$

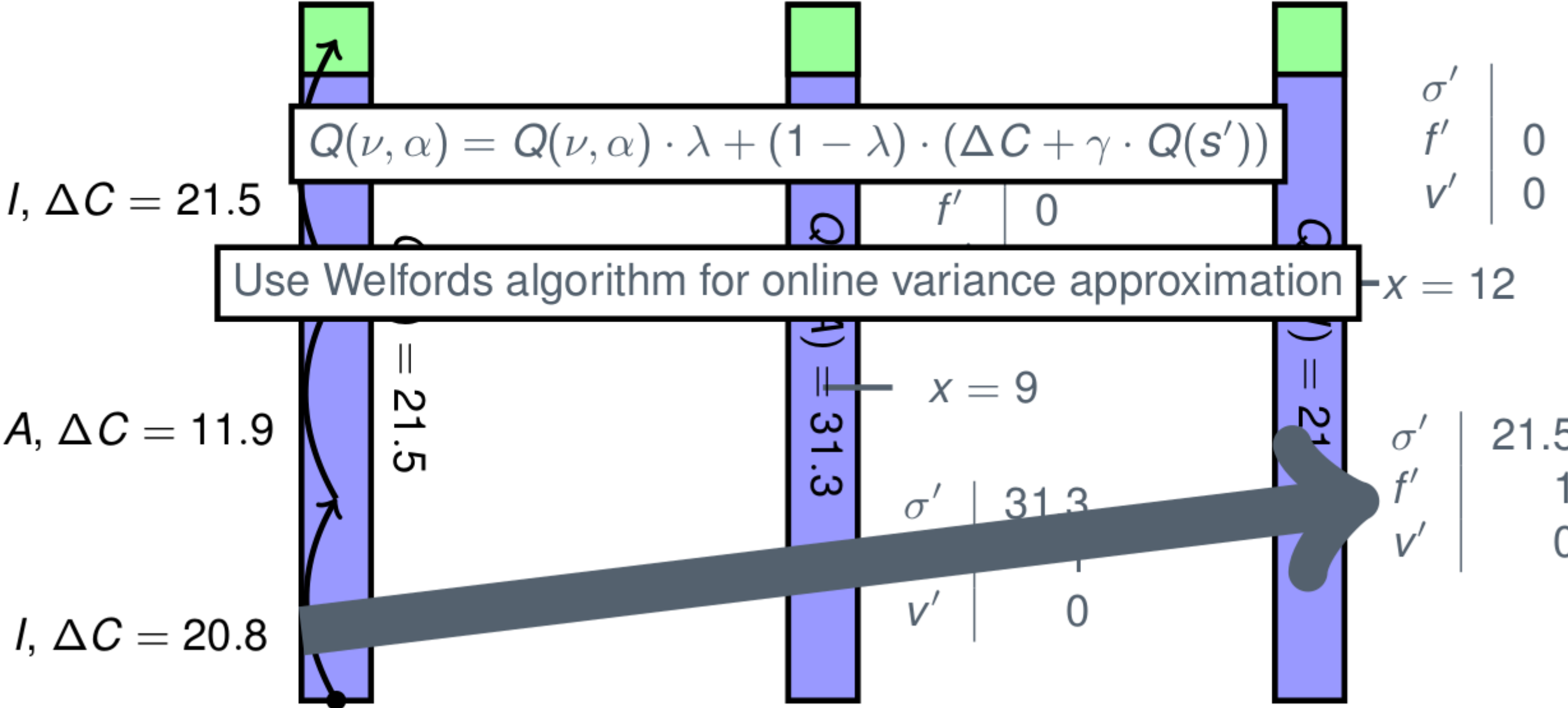
$$Q(\nu, I)$$



$$Q(\nu) = \min(Q(\nu, A), Q(\nu, I))$$

$$Q(\nu, A)$$

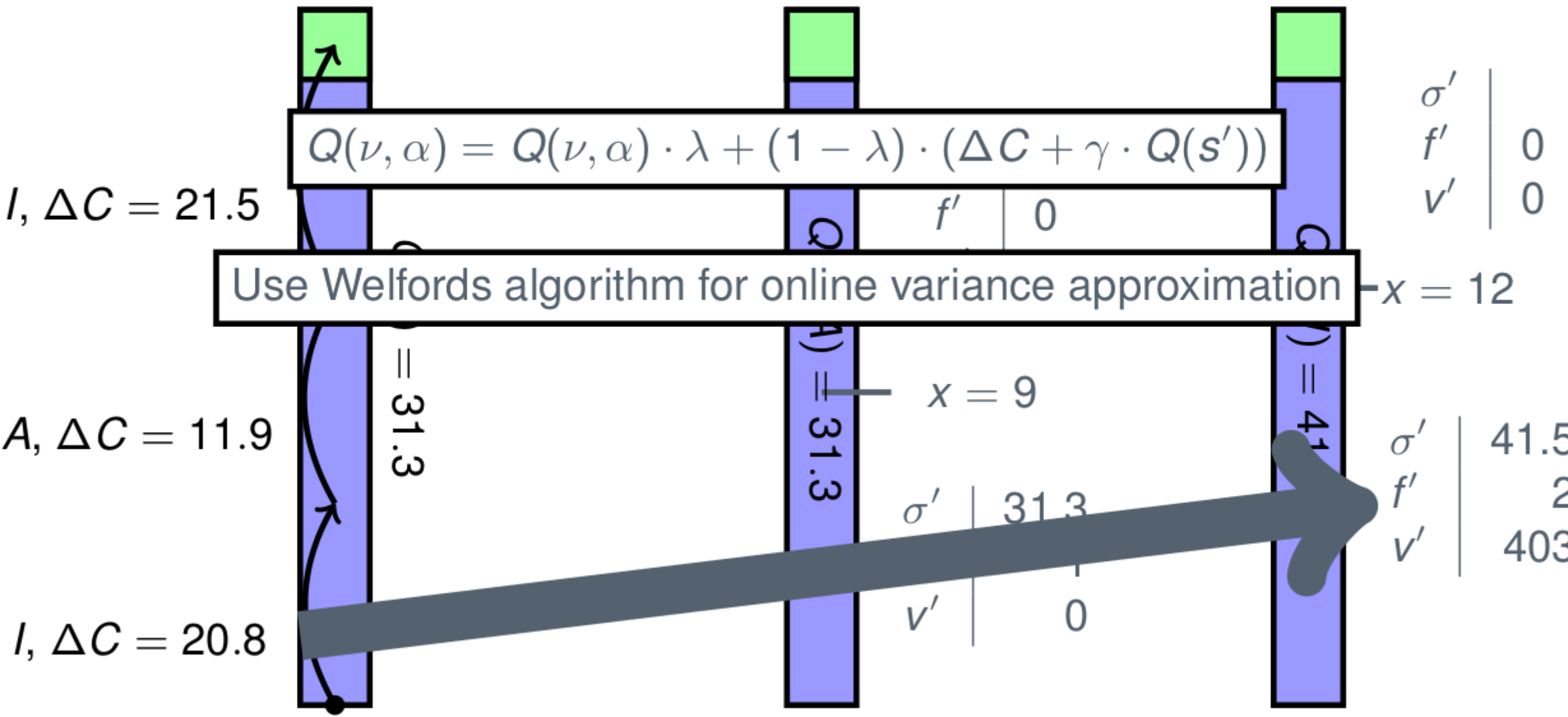
$$Q(\nu, I)$$



$$Q(\nu) = \min(Q(\nu, A), Q(\nu, I))$$

$$Q(\nu, A)$$

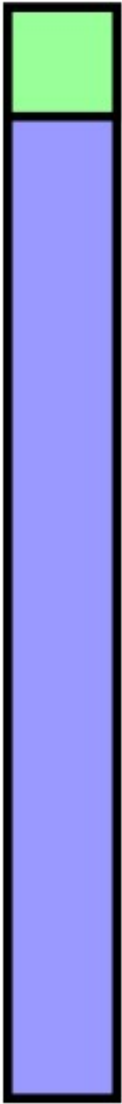
$$Q(\nu, I)$$



$$Q(\nu) = \min(Q(\nu, A), Q(\nu, I))$$

$$Q(\nu, A)$$

$$Q(\nu, I)$$



Sample some more



σ'	
f'	0
v'	0

σ'	31.3
f'	1
v'	0

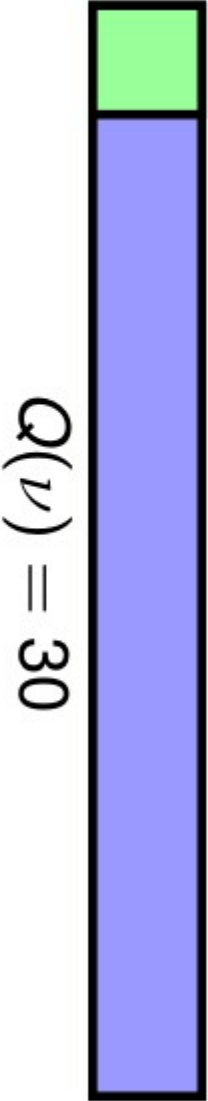


σ'	
f'	0
v'	0

$x = 12$

σ'	41.5
f'	2
v'	403

$$Q(\nu) = \min(Q(\nu, A), Q(\nu, I))$$



$$Q(\nu, A)$$



σ'	32.5
f'	15
v'	100

$$x = 9$$

σ'	34
f'	10
v'	900

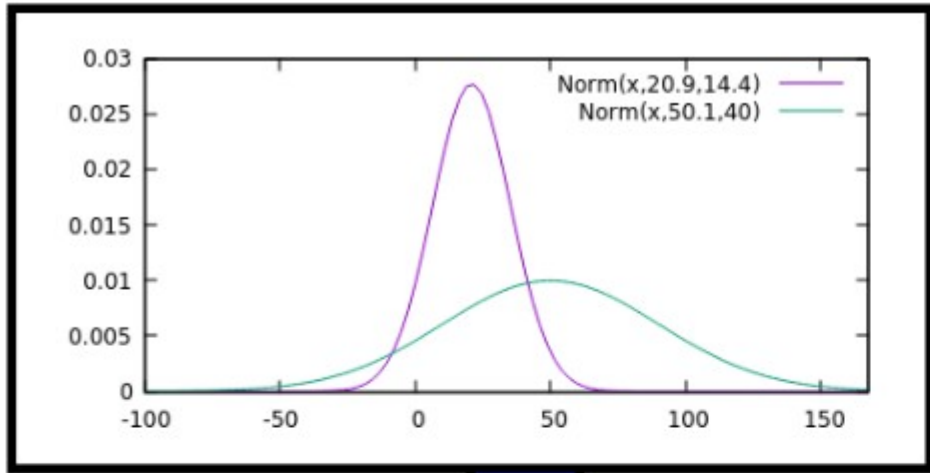
$$Q(\nu, I)$$



σ'	20.9
f'	19
v'	200

$$x = 12$$

σ'	50.1
f'	22
v'	1600

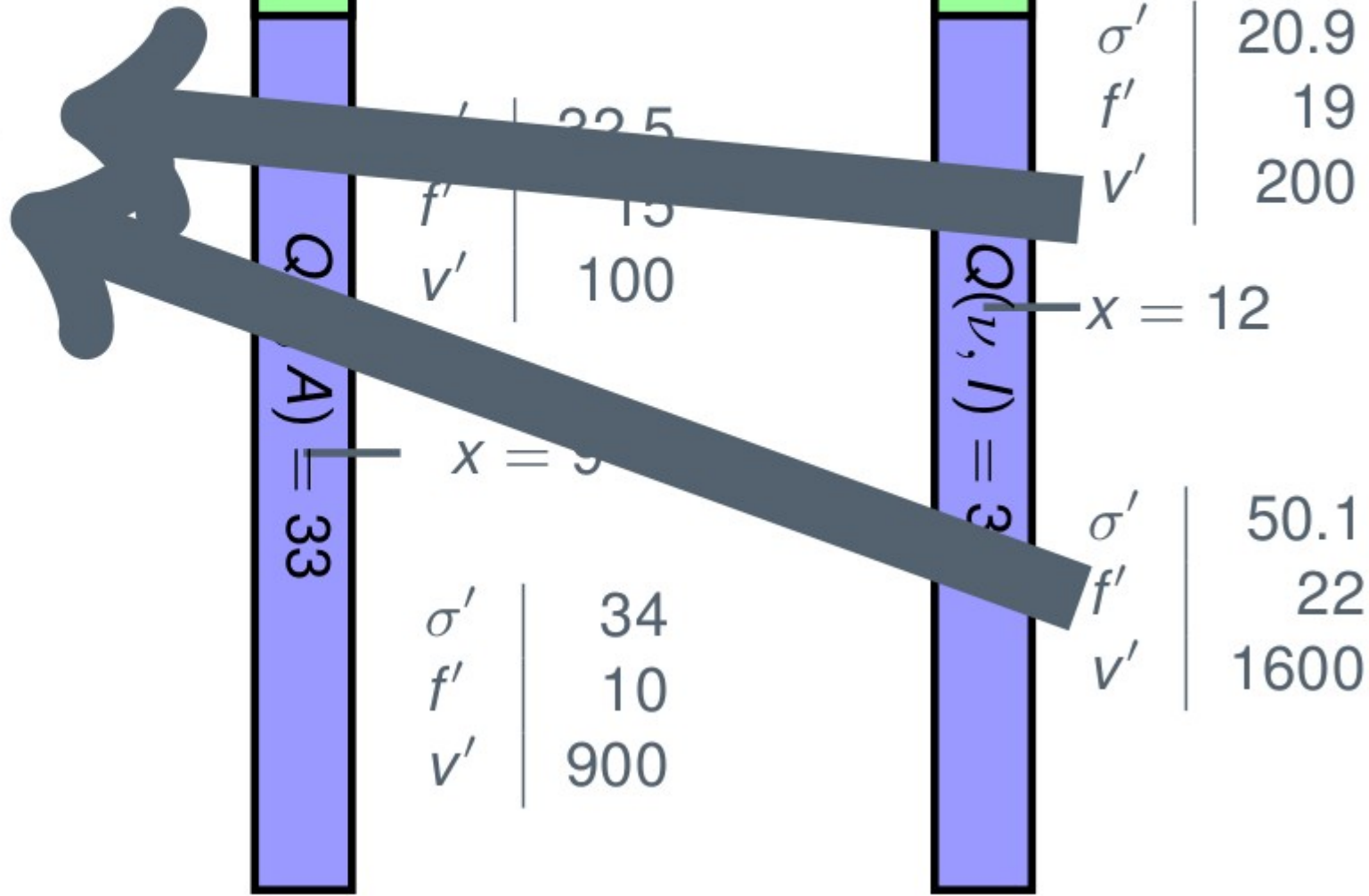


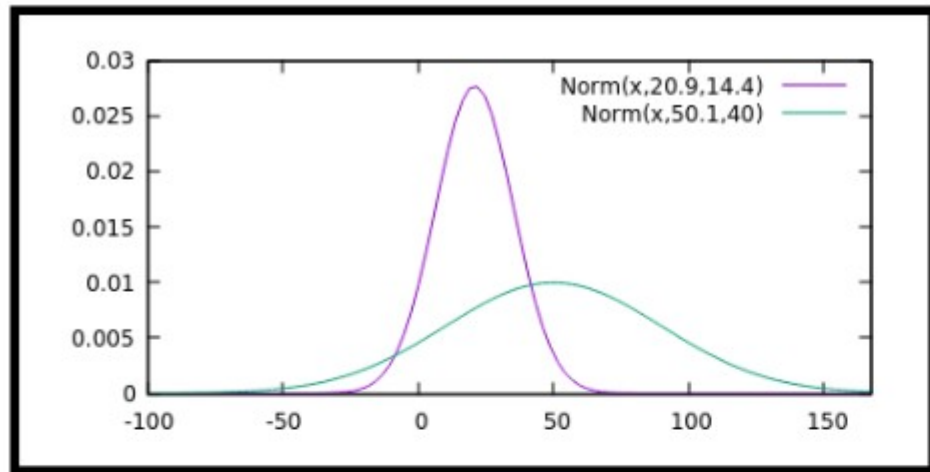
Welches T-test

$Q(\nu) = 30$

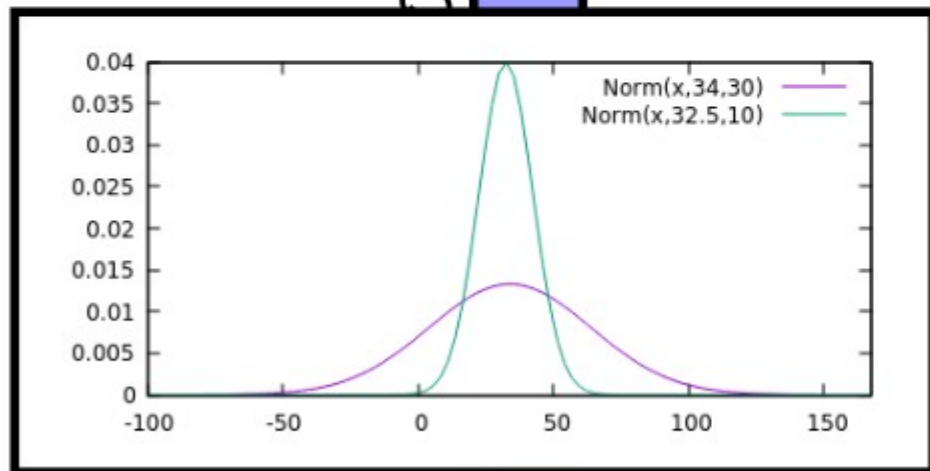
$Q(\nu, A)$

$Q(\nu, I)$



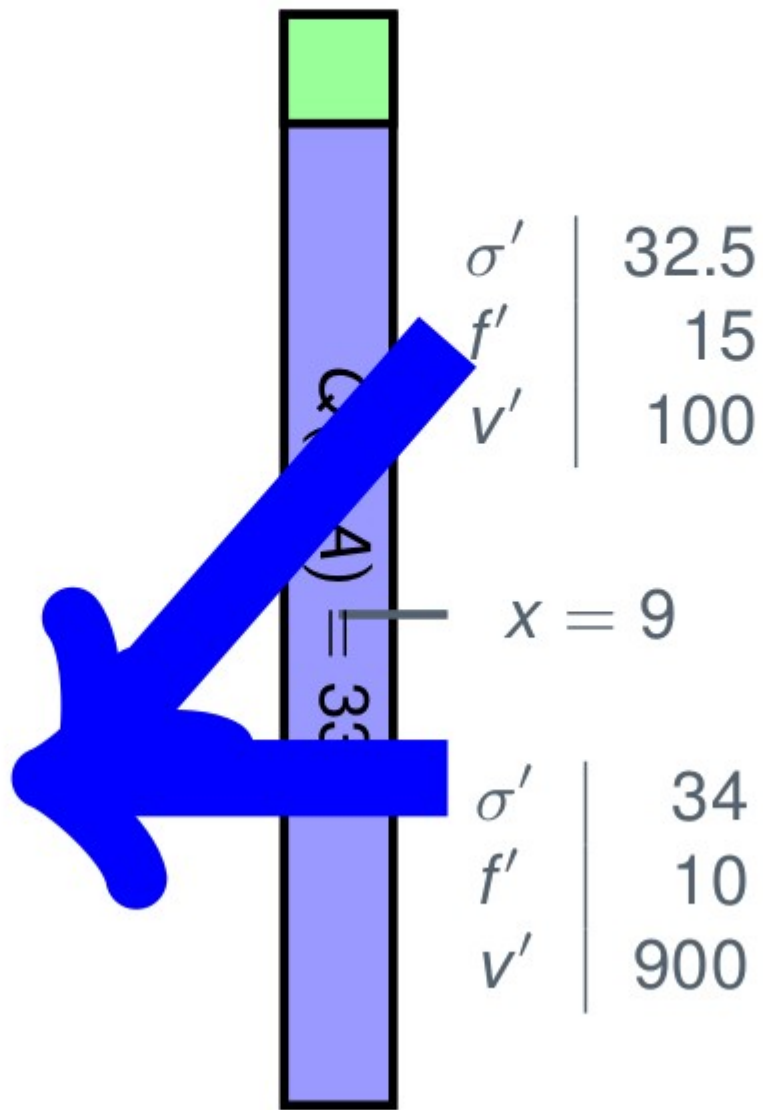


Welches T-test

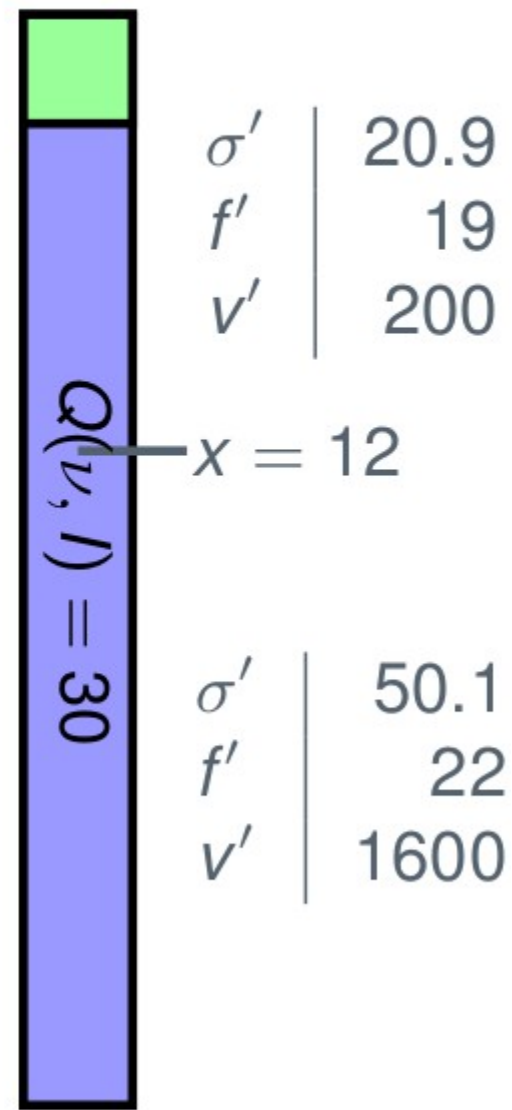


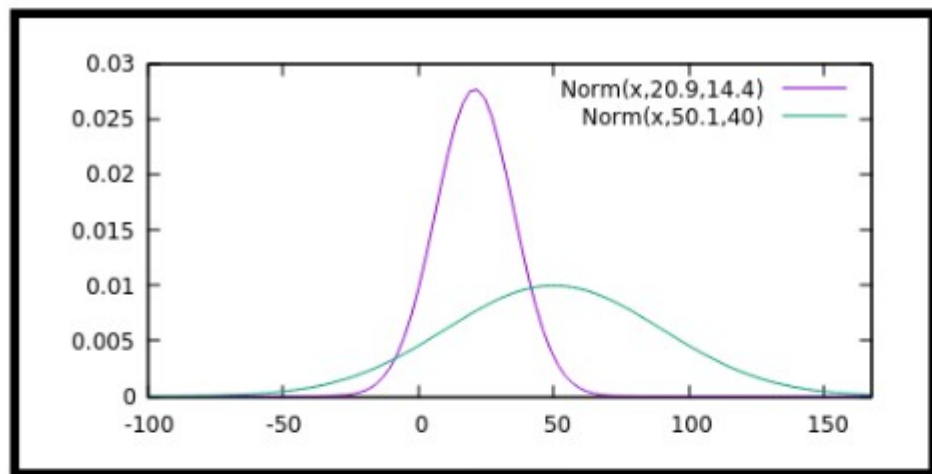
Kolmogorow-Smirnoff test

$Q(\nu, A)$

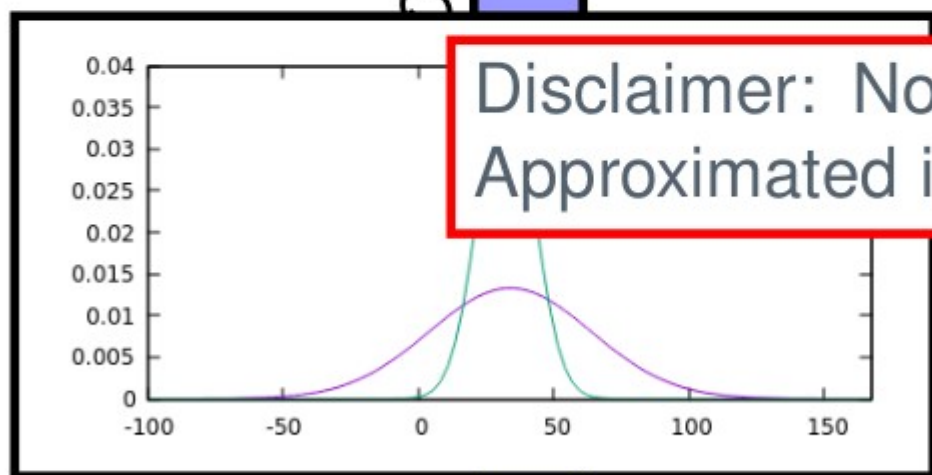


$Q(\nu, I)$





Welches T-test



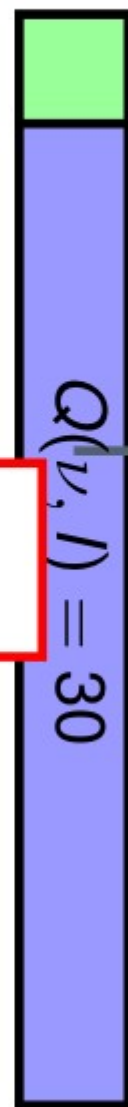
Kolmogorow-Smirnoff test

$Q(\nu, A)$



σ'	32.5
f'	15
v'	100

$Q(\nu, I)$



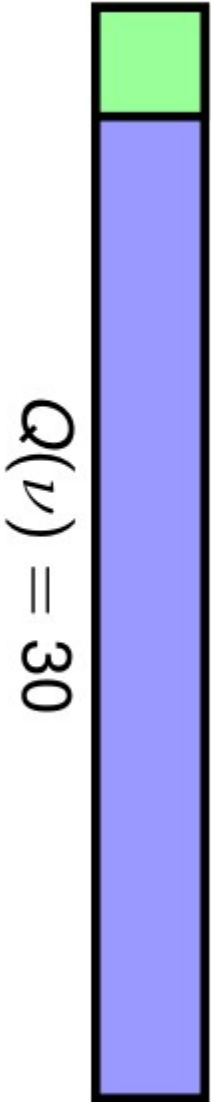
σ'	20.9
f'	19
v'	200

$x = 12$

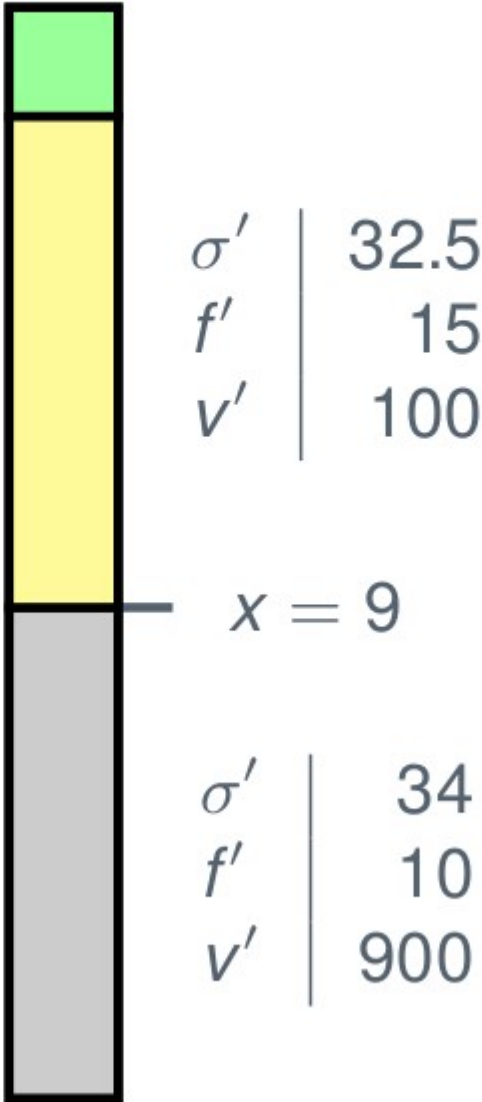
σ'	50.1
f'	22
v'	1600

Disclaimer: Nonsense statistics; it is a heuristic
Approximated in implementation

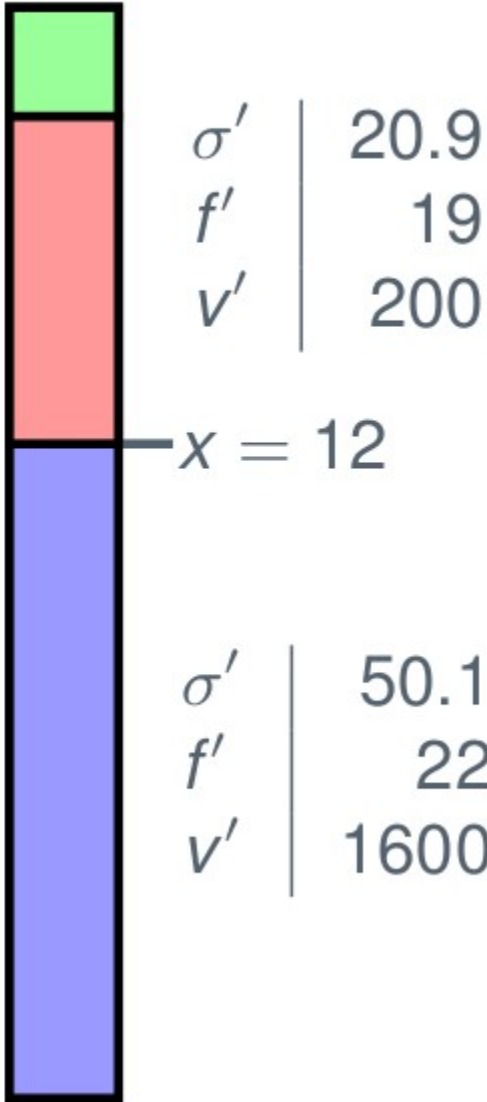
$$Q(\nu) = \min(Q(\nu, A), Q(\nu, I))$$



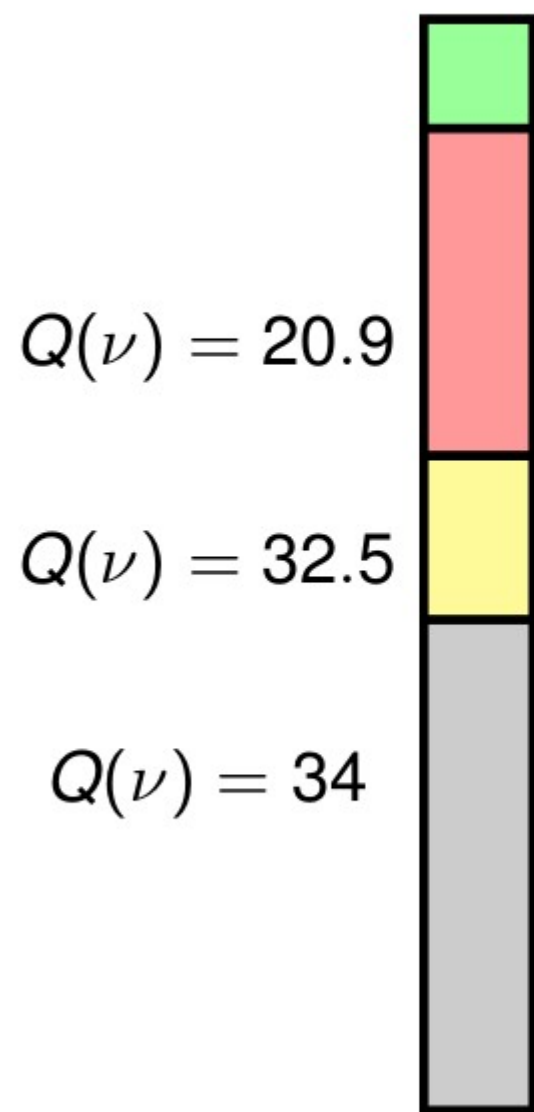
$$Q(\nu, A)$$



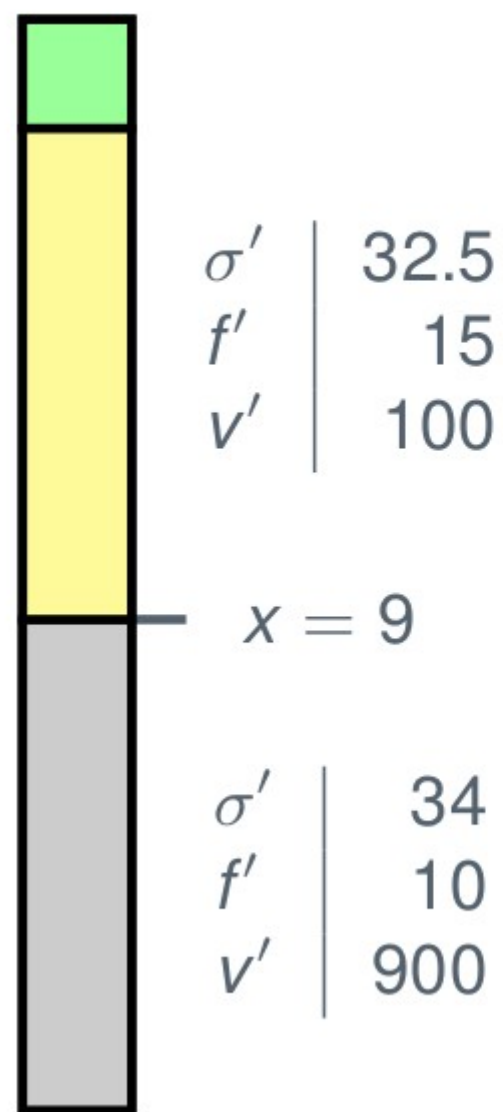
$$Q(\nu, I)$$



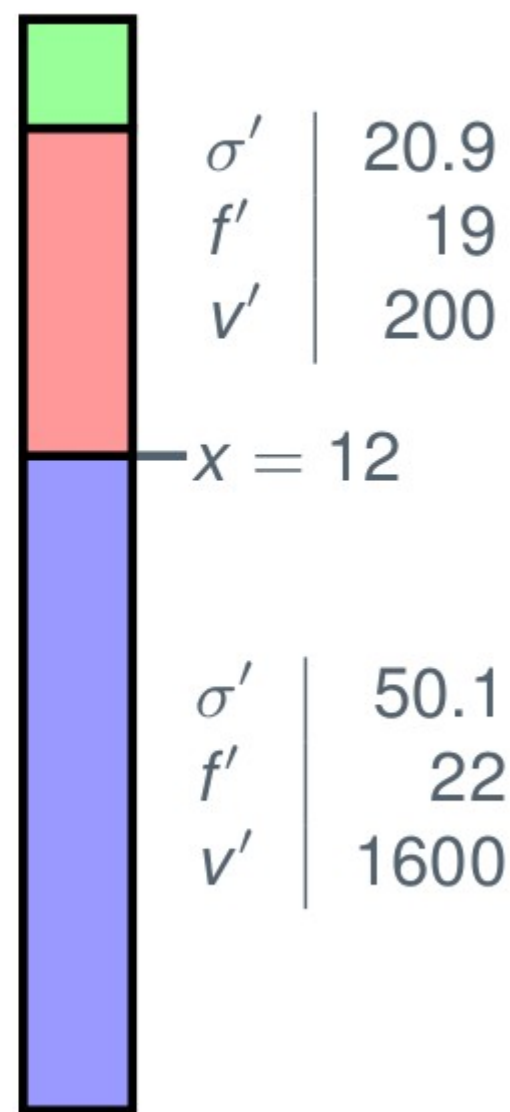
$$Q(\nu) = \min(Q(\nu, A), Q(\nu, I))$$



$$Q(\nu, A)$$



$$Q(\nu, I)$$



$$Q(\nu) = \min(Q(\nu, A), Q(\nu, I))$$

$$Q(\nu, A)$$

$$Q(\nu, I)$$

$$Q(\nu) = 20.9$$

$$Q(\nu) = 32.5$$

$$Q(\nu) = 34$$

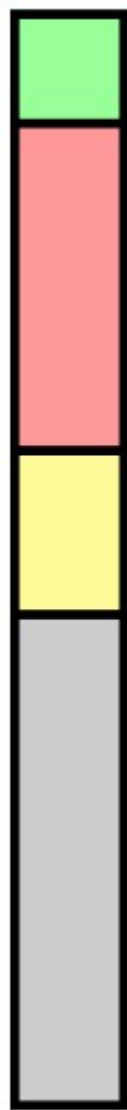
Sample more and repeat

$$Q(\nu) = \min(Q(\nu, A), Q(\nu, I))$$

$$Q(\nu) = 20.9$$

$$Q(\nu) = 32.5$$

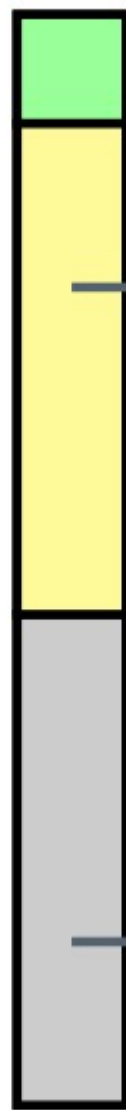
$$Q(\nu) = 34$$



$$Q(\nu, A)$$

$$x = 15$$

$$x = 3$$



$$Q(\nu, I)$$

$$x = 16$$

$$x = 6$$

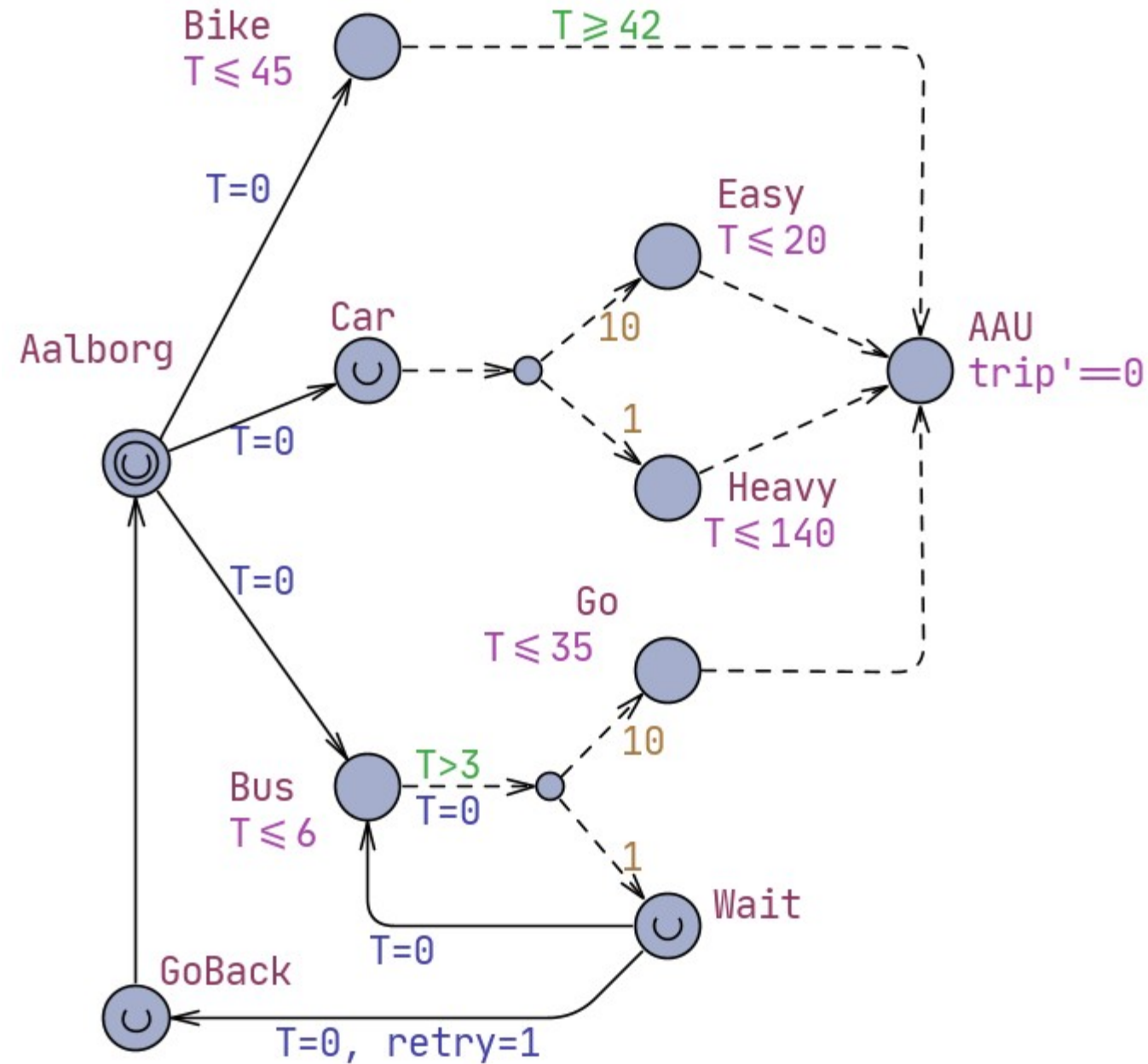


DEMO

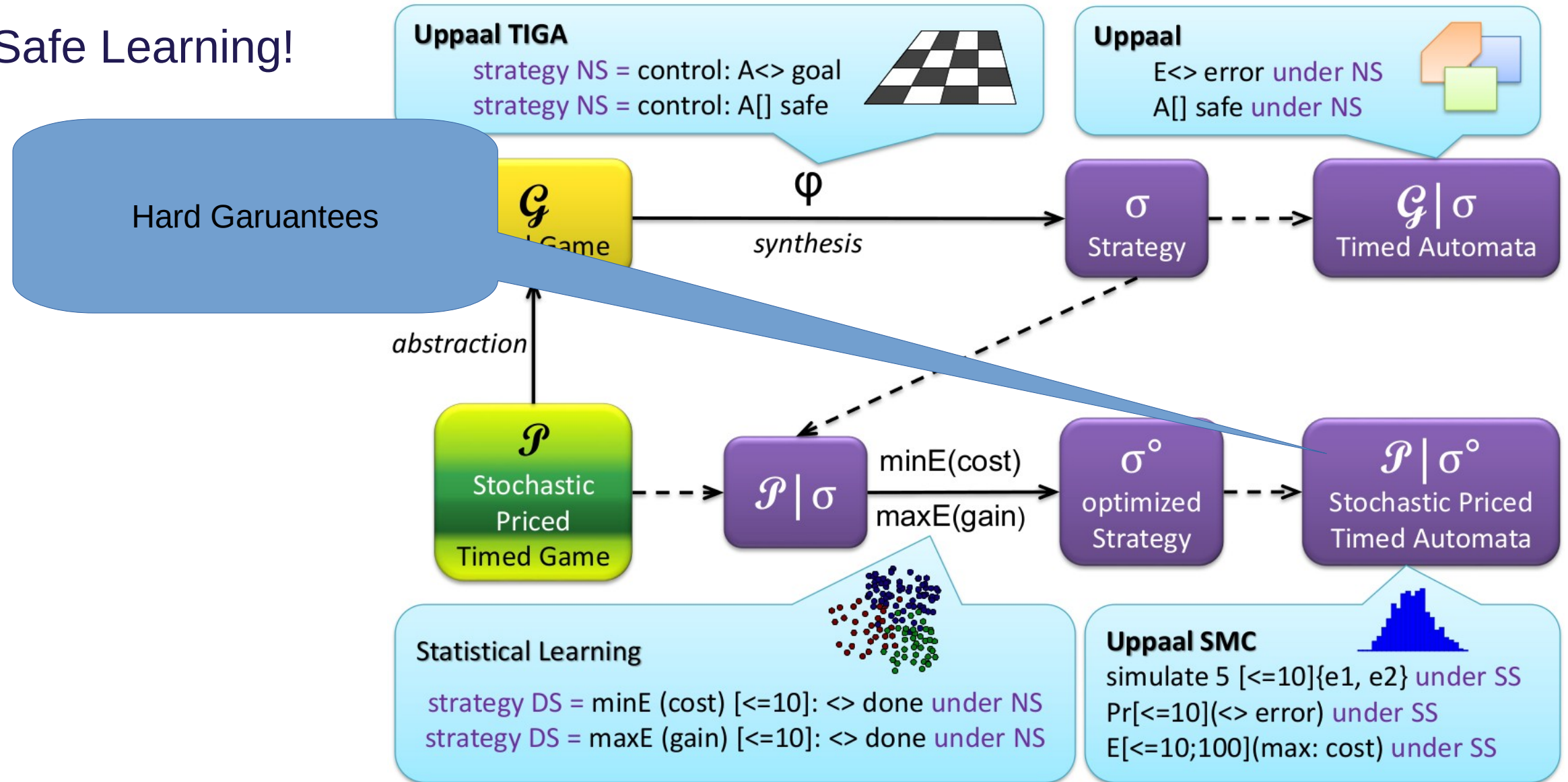
Safe Learning!

The curious case of going to work

What is the fastest way to get to work?
What if Kim needs to be there in 40 min?

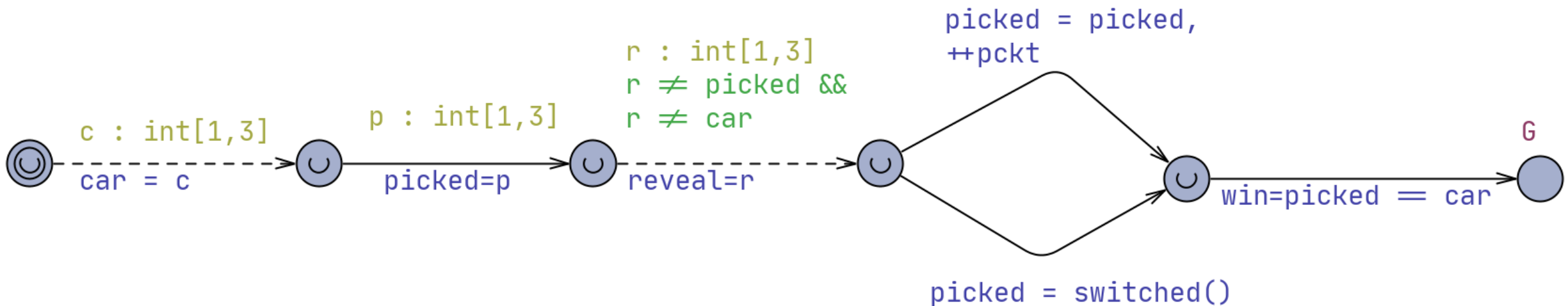


Safe Learning!



Partial Observability in Stratego

The Monty Hall Problem



We cannot observe what is behind the door!

```
strategy s = maxE(win) [<=1] : <> Hall.G // observes everything
```

```
strategy s = maxE(win) [<=1] {Hall.location, reveal} -> {}: <> Hall.G
```

Partial Observability in Stratego

The query explained

strategy s = maxE(...) [≤...] { ... } -> { ... } : <> ...

Timebound
e.g. time ≤ 100

Goal-condition, often
inverse of timebound:
time ≥ 100 || WIN

Cost expression

Variables to be explicit about
I.e. variables with few values
Locations, booleans ...
Defaults to all locations, ints, bools

Variables to generalize over
I.e. variables with many values:
Doubles, clocks ...
Defaults to all clocks & doubles

Tips when using Stratego

Non-lazy controllers

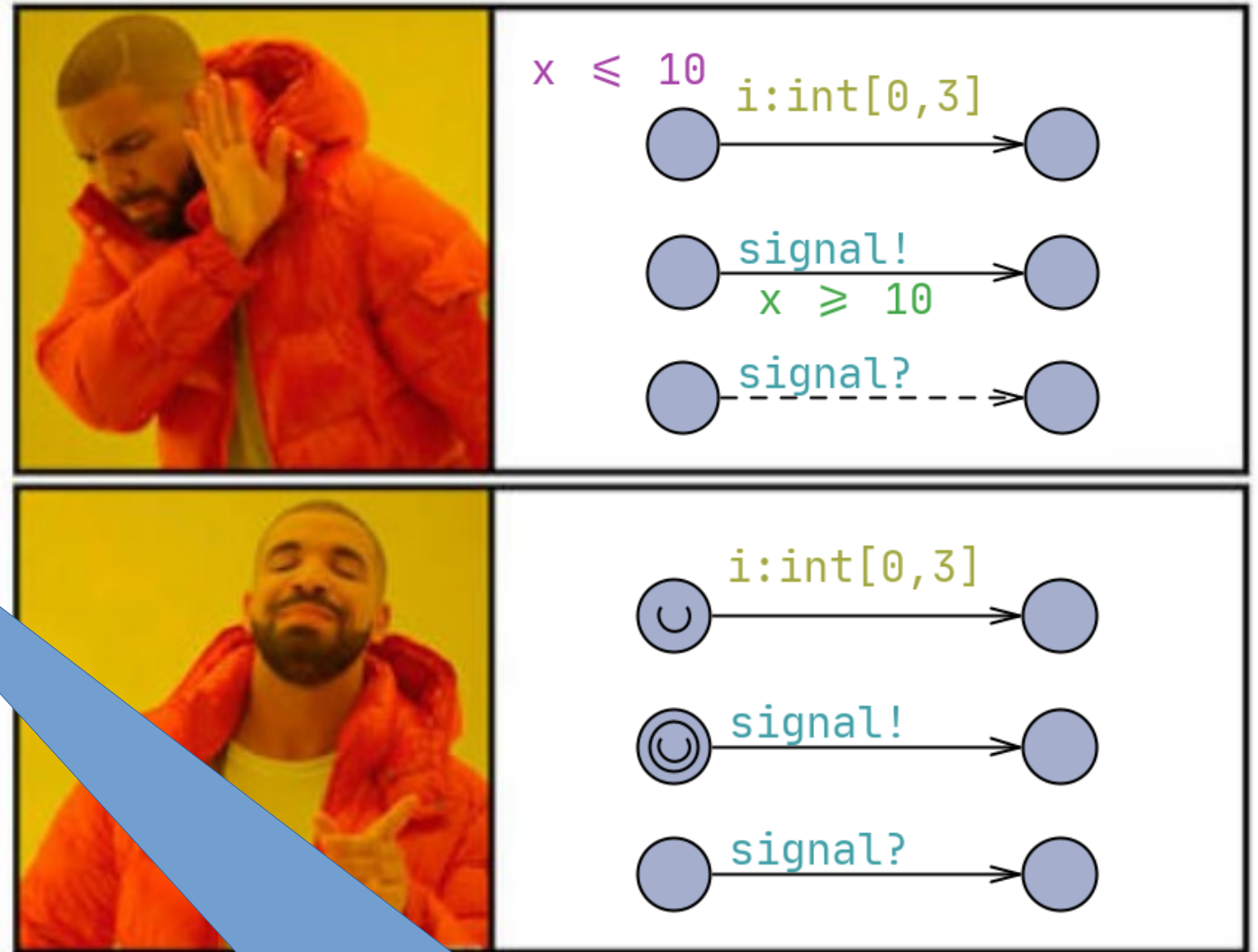
- Act now or “wait forever”
- Stratego cannot propose a delay!
- Keep “controllable locations” urgent
 - Avoid “lazy” behaviour

Problem is Markovian

- Future is independent from past
- Cost-function is Markovian
- The observable projection is Markovian!

Other Tips

- Keep receivers controllable
 - Avoids consistency-issue with TiGa
- Avoid guards on controllable edges
 - Breaks assumption of Q-learning

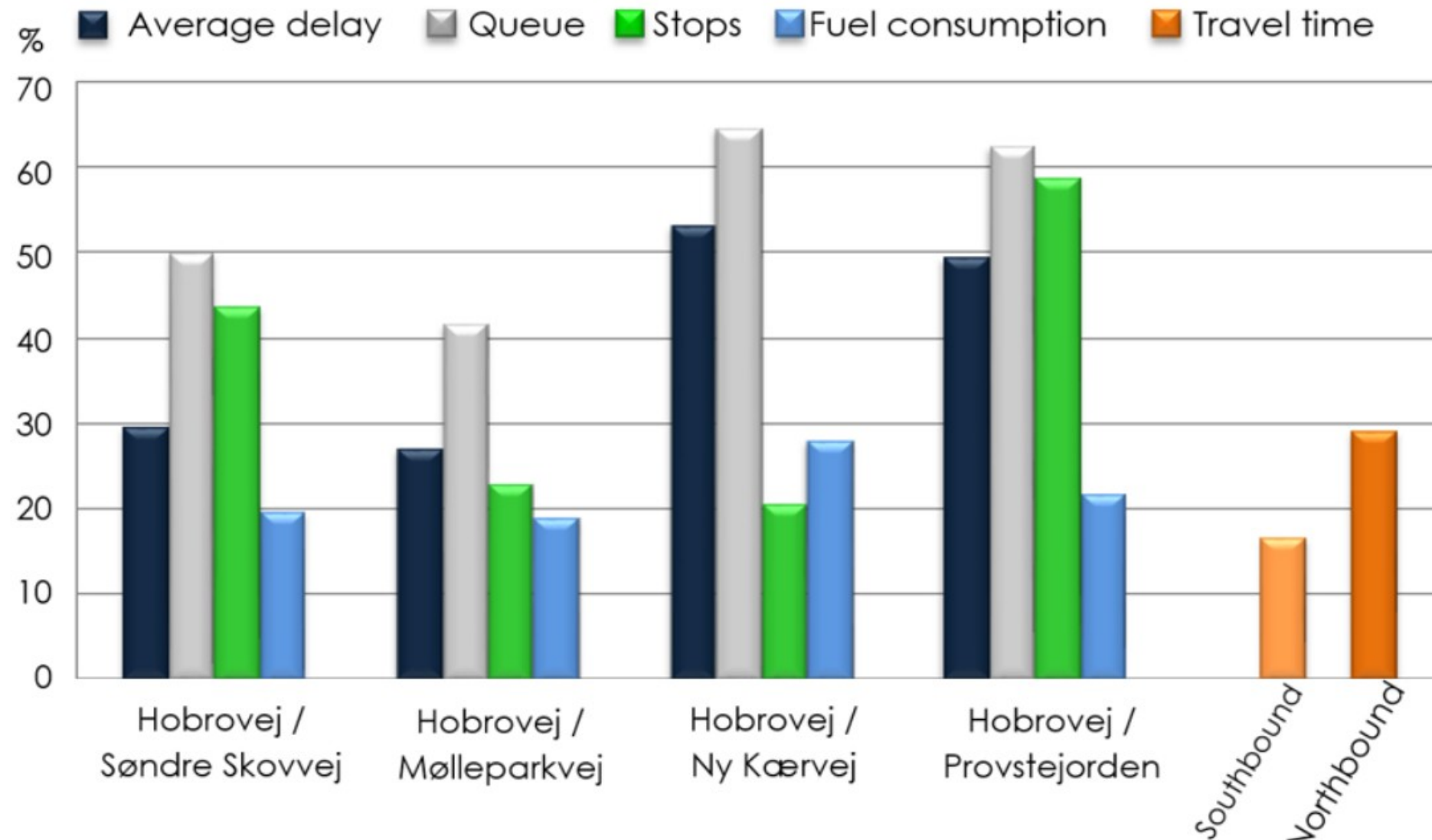


WIP to improve all of these things

Exercise!

3) Stochastic Jobshop Scheduling Game

- Extend your Stochastic Jobshop scheduling to a stochastic game
- Optimize the scheduling using Stratego, minimize the completion time
- Make a strategy (using TiGa) that guarantees that Kim is done in 60 minutes
 - He has a plane to catch!
- Compare the random scheduler, the fast scheduler and the “Kim gets to the plane”-scheduler
- Try to improve (using Stratego) the “Kim gets to the plane”-strategy
- Change the observations of stratego s.t. only the “sections in use” and the location of the persons are visible, what is the effect?



AALBORG
KOMMUNE



AARHUS
KOMMUNE

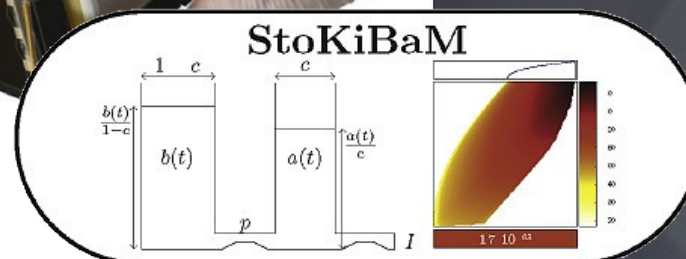
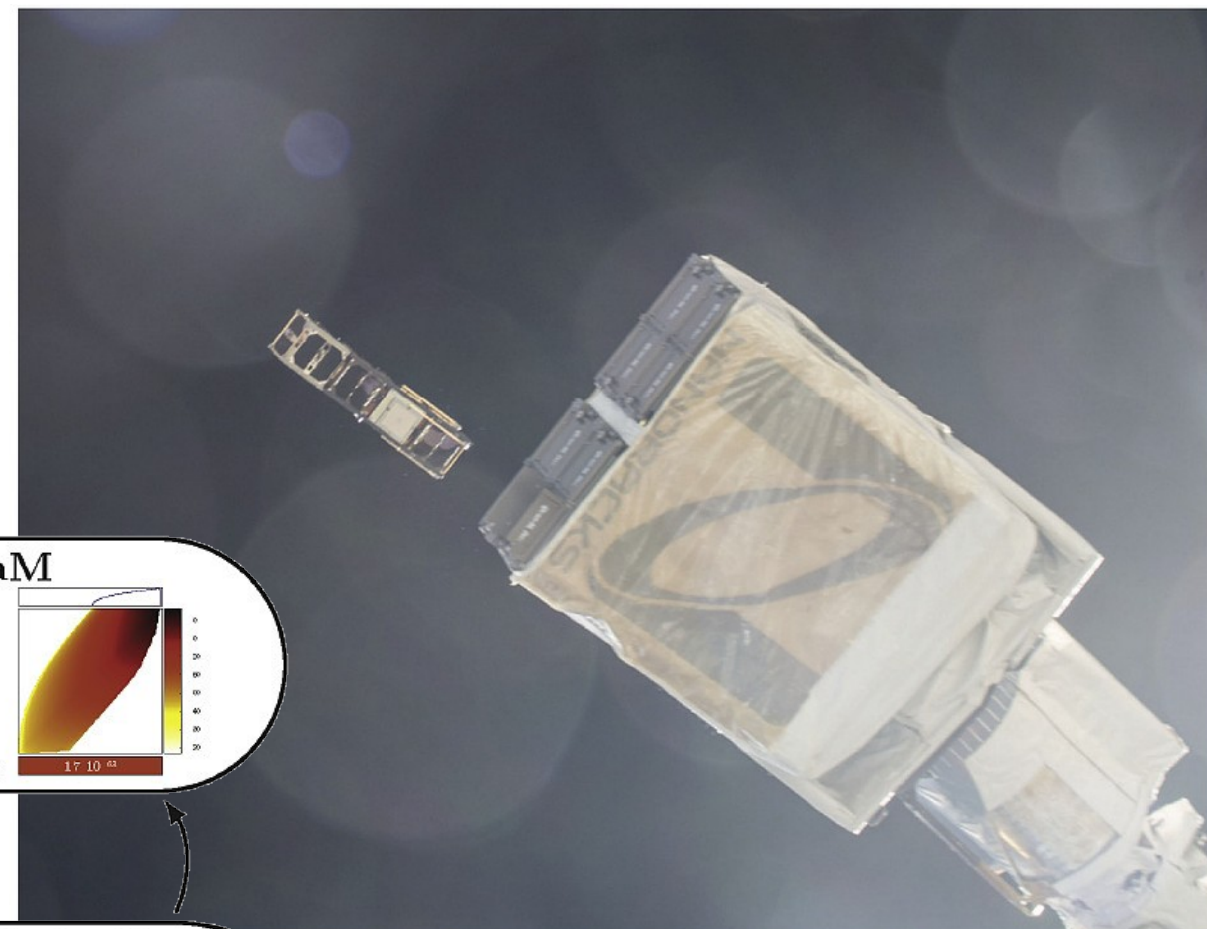


KØBENHAVNS KOMMUNE

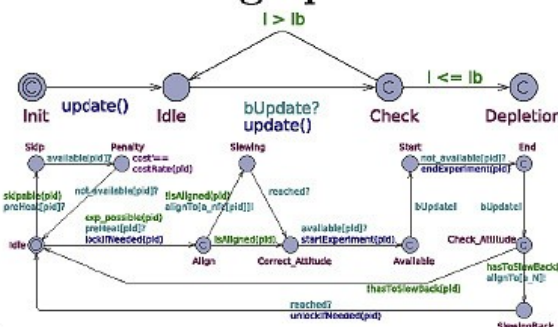
COWI

RAMBOLL

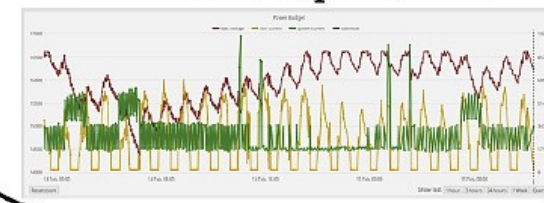




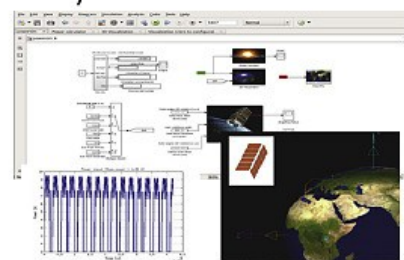
Scheduling optimization



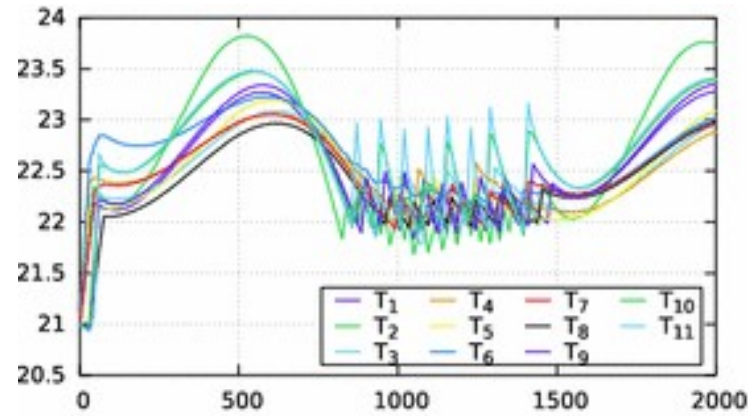
Satellite Operation



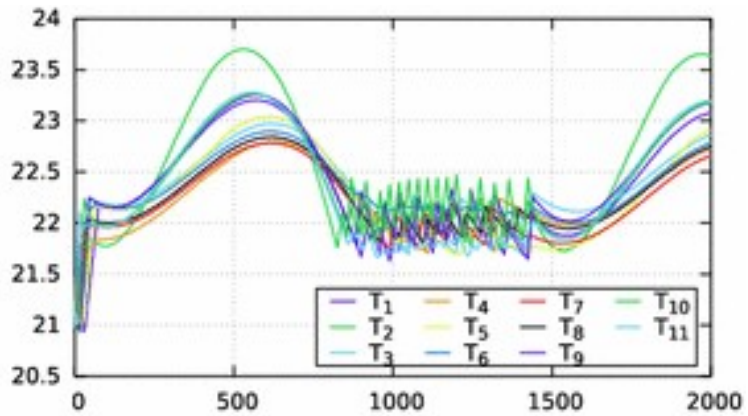
Orbit/Power Prediction



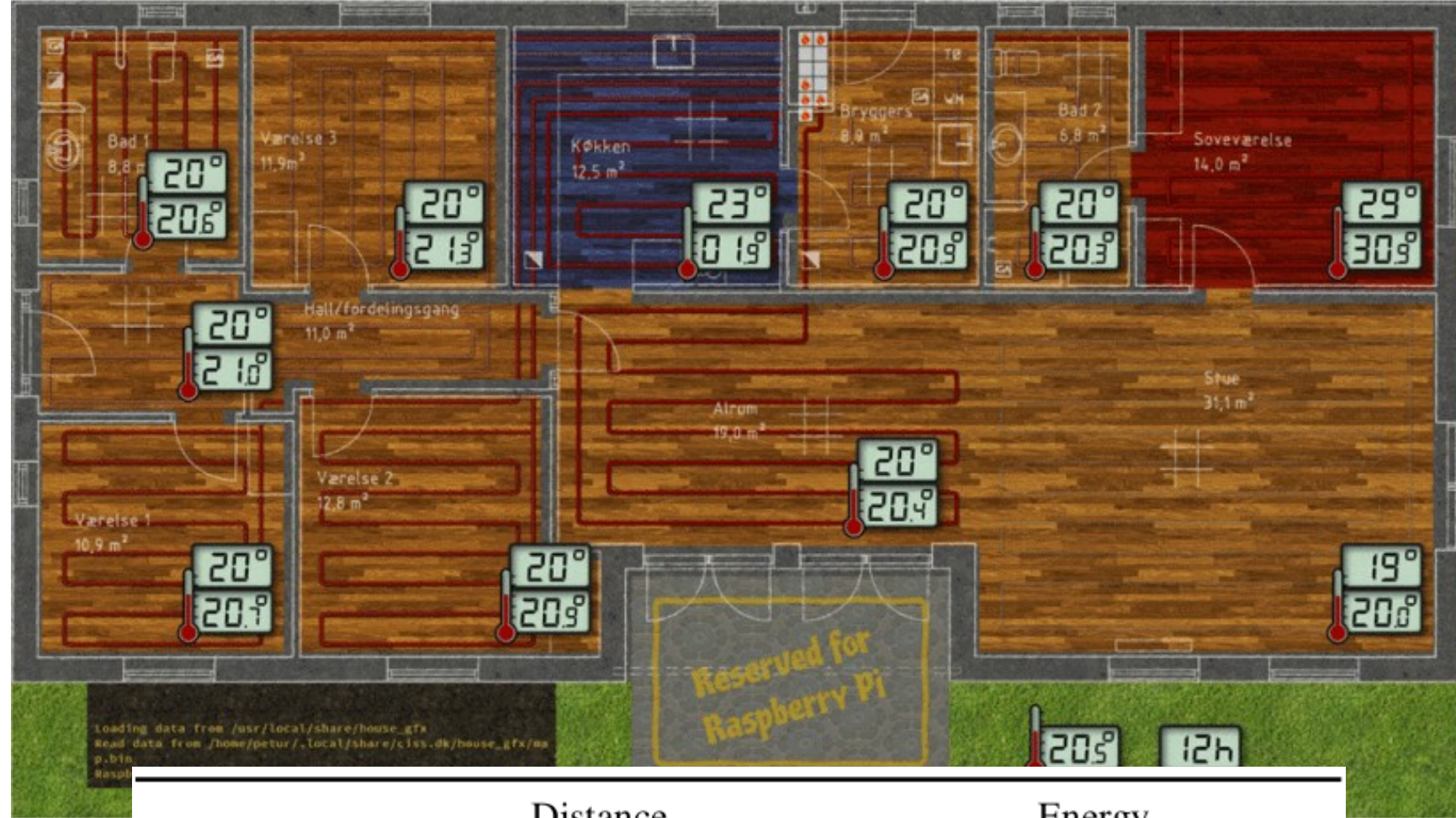
Examples



(a) Bang-Bang controller



(b) STRATEGO-ON-CL controller



Weather	Distance			Energy		
	Bang-Bang	STRATEGO	imp.	Bang-Bang	STRATEGO	imp.
Aalborg	14583	8342	43%	14180	12626	10%
Anadyr	2385515	1483272	37%	23040	22475	2%
Ankara	17985	10464	41%	17468	15684	10%
Minneapolis	22052	12175	44%	18165	15882	12%
Murmansk	399421	187941	52%	22355	21011	6%

Things we have not discussed

- Differential equations
- Loading/saving strategies
- Interfacing with external C-code
- Loading real-life data
- On-line reinforcement learning
 - Adaptive/data-driven planning