

COGSCI 2026

Effective Explanations Support Planning Under Uncertainty

Hanqi Zhou · Britt Besch · Charley M. Wu · Tobias Gerstenberg

University of Tübingen · TU Darmstadt · MPI for Biological Cybernetics

hanqi.zhou@uni-tuebingen.de



 GitHub: [cicl-stanford/
explaining_how_cogsci26](https://github.com/cicl-stanford/explaining_how_cogsci26)

How to get food during the lunch break?

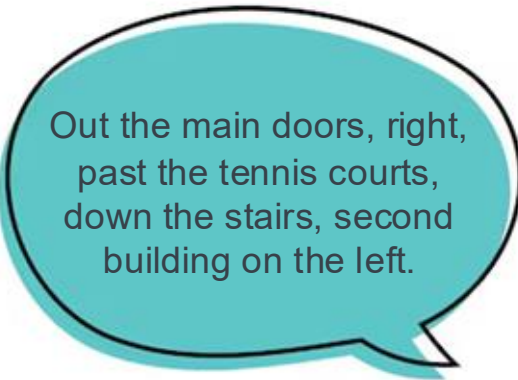


There's amazing
food in Leblon.

How to get food during the lunch break?



There's amazing
food in Leblon.



Out the main doors, right,
past the tennis courts,
down the stairs, second
building on the left.

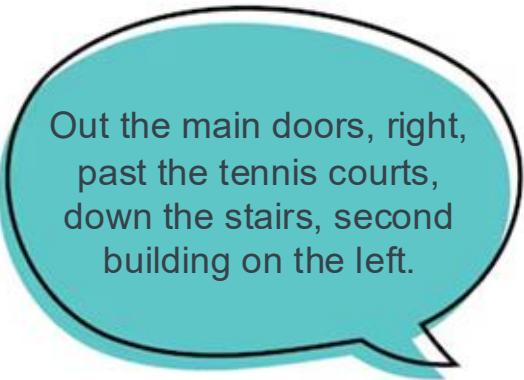
How to get food during the lunch break?



There's amazing
food in Leblon.



There are a couple
of places by the
pool, ground floor.



Out the main doors, right,
past the tennis courts,
down the stairs, second
building on the left.

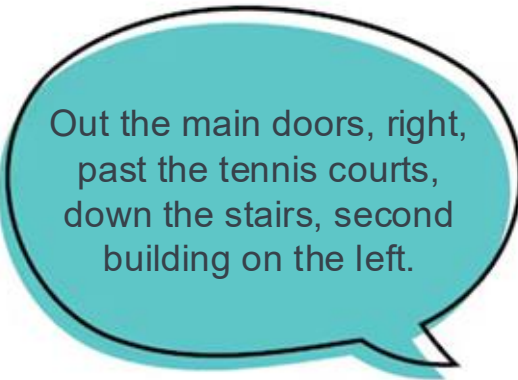
How to get food during the lunch break?



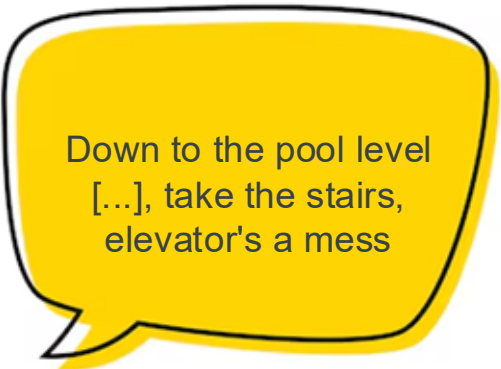
There's amazing
food in Leblon.



There are a couple
of places by the
pool, ground floor.



Out the main doors, right,
past the tennis courts,
down the stairs, second
building on the left.



Down to the pool level
[...], take the stairs,
elevator's a mess

THE QUESTION



- LIMITS: attention, memory, words, etc.
- ASYMMETRY: the explainer knows the task and the environment; the listener doesn't
- UNCERTAINTY: the explainer doesn't know what the listener will be facing when they use it

- LIMITS: attention, memory, words, etc.
- ASYMMETRY: the explainer knows the task and the environment; the listener doesn't
- UNCERTAINTY: the explainer doesn't know what the listener will be facing when they use it

A good explanation selects what matters and structures it so the listener can act.

From belief to action

$$S(e|w) \propto \exp(\lambda \cdot U(e, w))$$

utility = **help to the listener** – cost to say it

MOST ACCOUNTS

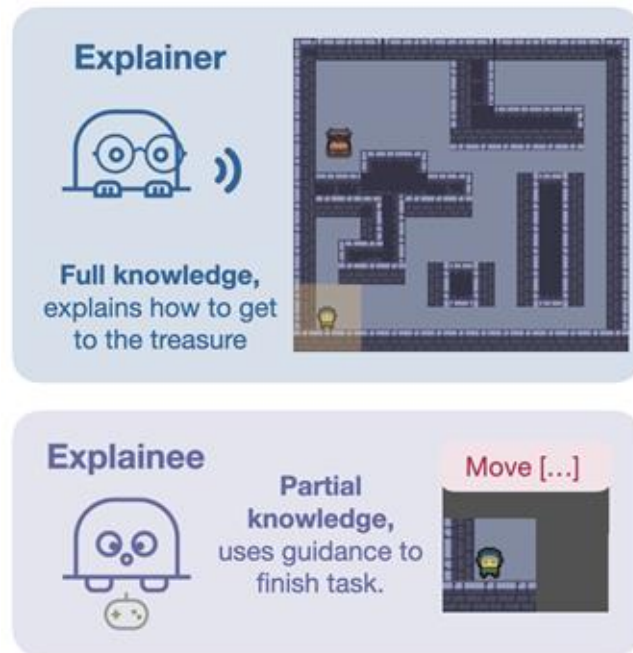
help = right **belief**

OURS

help = right **action** under the updated belief

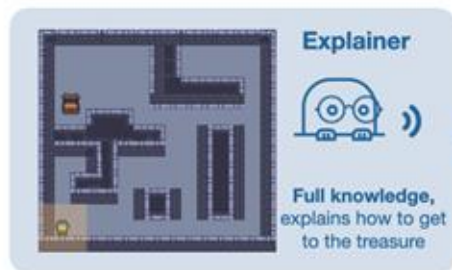
A controlled navigation

- Explainer sees the FULL map
- Navigator sees only a LOCAL window
- One free-text message, one shot

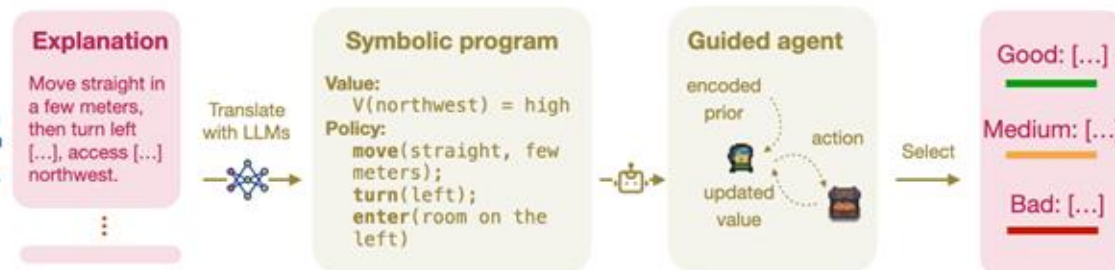


Make a model act on the explanation

Explanation Collection



Explanation Selection

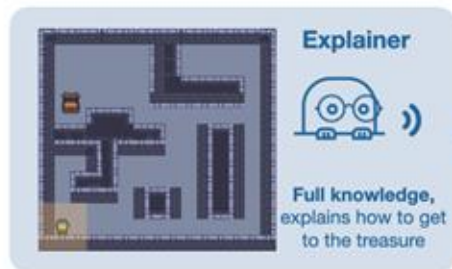


Collect

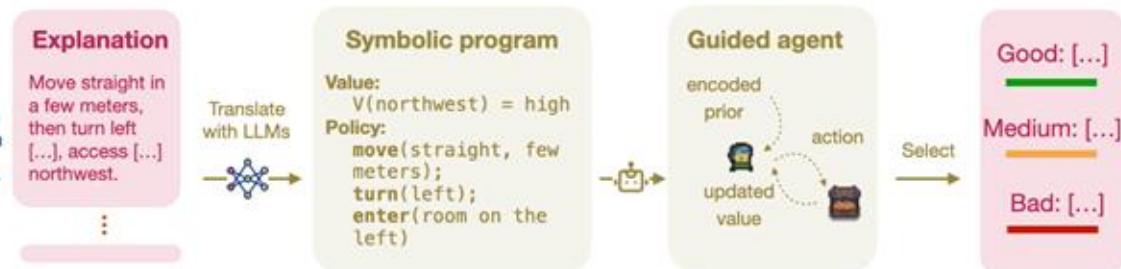
real human explanations,
written from the full map

Make a model act on the explanation

Explanation Collection



Explanation Selection

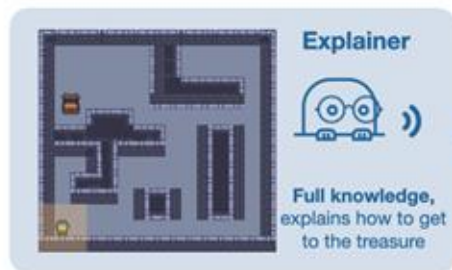


Compile

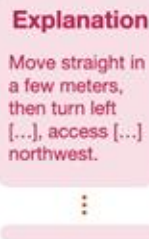
an LLM turns each message into a small runnable program

Make a model act on the explanation

Explanation Collection



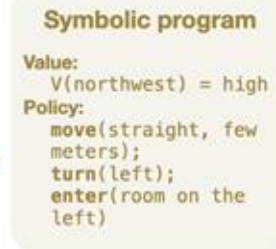
①
Provide an
explanation



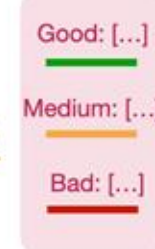
Translate
with LLMs



Explanation Selection



Select

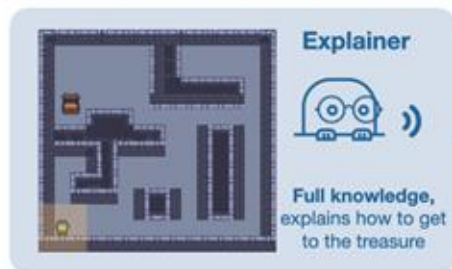


Plan & act

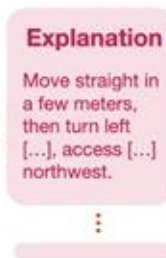
an agent runs the program from a local view

Make a model act on the explanation

Explanation Collection



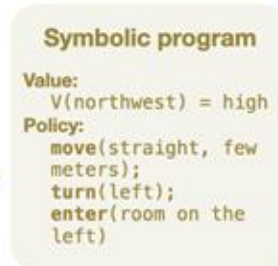
①
Provide an
explanation



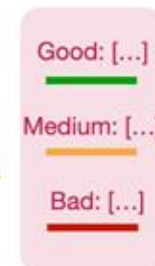
Translate
with LLMs



Explanation Selection



Select

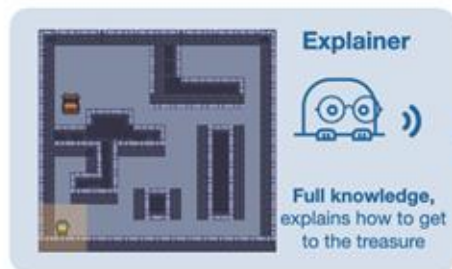


Score

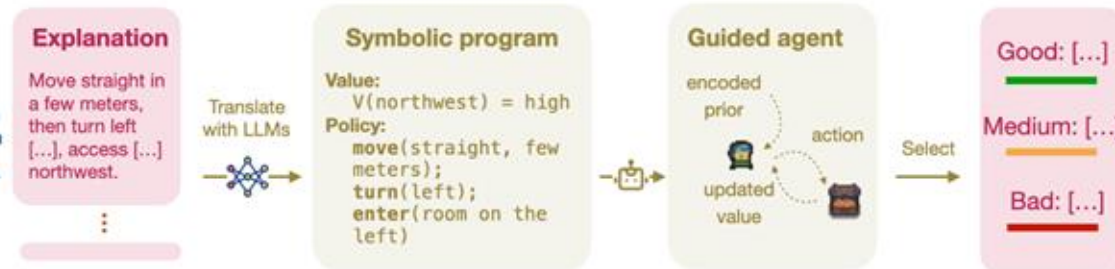
rate the resulting
behaviour → good
/ medium / bad

Make a model act on the explanation

Explanation Collection



Explanation Selection



Compile

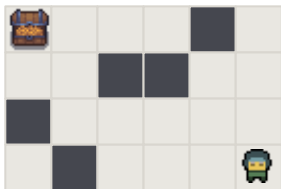
an LLM turns each message into a small runnable program

Score

rate the resulting behaviour → good / medium / bad

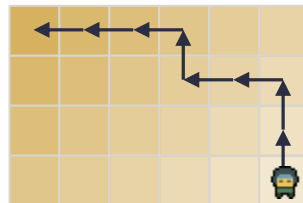
Component: compile

Actual world w &
Explanation e



Conveyed world (V, ρ)

← from e



V value prior (heat)

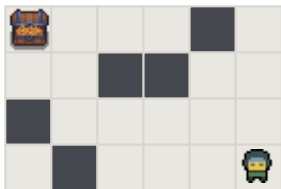
where reward is, spread over the whole map.

ρ policy prior (arrow)

a cell-by-cell hint of what to do.

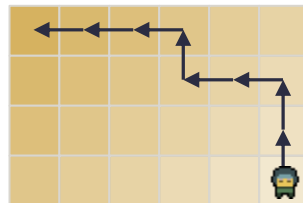
Component: compile

Actual world w &
Explanation e



Conveyed world (V, ρ)

← from e



V value prior (heat)

where reward is, spread over the whole map.

ρ policy prior (arrow)

a cell-by-cell hint of what to do.



“Go straight, turn right at the third opening, then keep going until you reach the wall. You’ll see the treasure there.”



Policy:

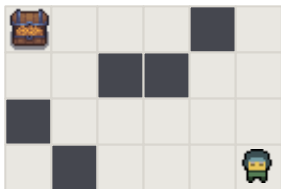
```
move_until(forward, opening==3)
turn(right)
move_until(forward, wall)
```

Value:

$V(\text{wall ahead}) = \text{high}$

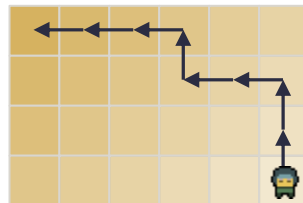
Component: compile

Actual world w &
Explanation e



Conveyed world (V, ρ)

← from e



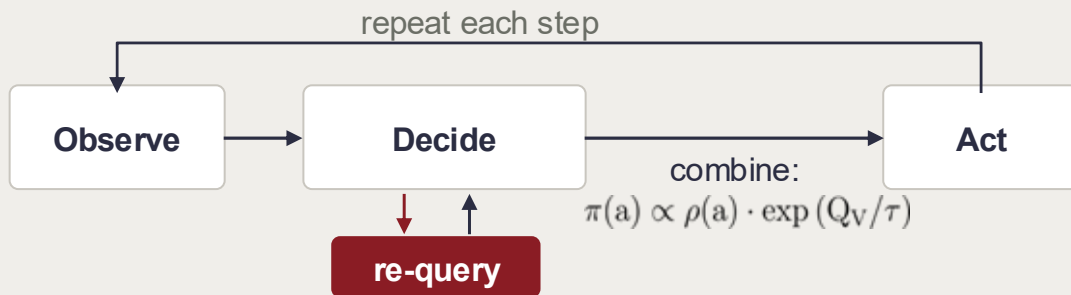
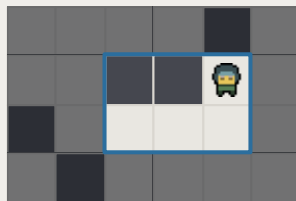
V value prior (heat)

where reward is, spread over the whole map.

ρ policy prior (arrow)

a cell-by-cell hint of what to do.

Plan $\pi_{(V, \rho)}$ the agent acts from a local view



priors run out / dead-end → planning cost

Component: utility



AN EXPLANATION

“Go straight, turn right at the third opening, then keep going until you reach the wall. You’ll see the treasure there.”

ITS CONVEYED WORLD

Policy:

```
move_until(forward, opening==3)
turn(right)
move_until(forward, wall)
```

Value:

$V(\text{wall ahead}) = \text{high}$

Component: utility



AN EXPLANATION

“Go straight, turn right at the third opening, then keep going until you reach the wall. You’ll see the treasure there.”

ITS CONVEYED WORLD

Policy:

```
move_until(forward, opening==3)
turn(right)
move_until(forward, wall)
```

Value:

$V(\text{wall ahead}) = \text{high}$

WHAT WE MEASURE

Did it resolve into action?

≈ 0.4

REPLAN · re-queries when stuck · lower

Was it efficient?

$1.1\times$

LEN · realized vs. shortest path · lower

Did it arrive?

95%

SUCC · reached the goal in budget · higher

$R(\pi^*) = \gamma \cdot \text{SUCC} - \alpha \cdot \text{REPLAN} - \beta \cdot \text{LEN} \rightarrow \text{Good}$

Component: utility

STANDARD RSA

$$S(e|w) \propto \exp(\lambda \cdot U(e, w))$$

$$U(e, w) = E_{\hat{w} \sim L(\hat{w}|e)} [\log L_0(w|e)] - C(e)$$

Help: listener assigns probability to the intended world

OURS

$$S(e|w) \propto \exp(\lambda \cdot U(e, w))$$

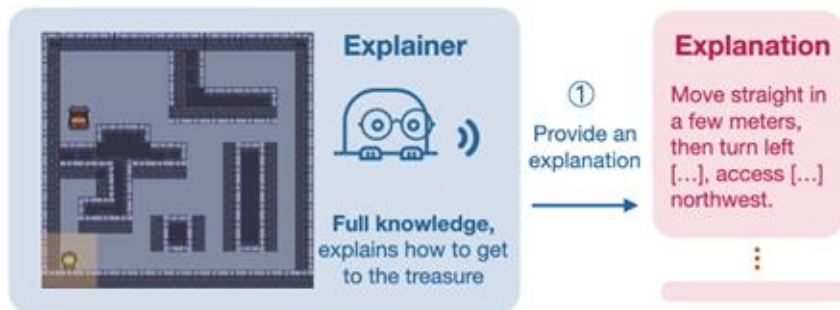
$$U(e, w) = E_{z \sim q_0(z|e), \tau \sim P(\tau|w, \pi_z)} [r(\tau, w)] - C(e)$$

help: listener's action produces reward in the actual world

$z=(V, \rho)$: V = value prior, ρ = policy prior
 τ actual trajectory

What people say · think · do

4 preregistered experiments · 1,200 explanations · 24 maps · run on Prolific



1

Collection

N = 50

Explainers see the full map and write a message to guide a partner.

What people say · think · do

4 preregistered experiments · 1,200 explanations · 24 maps · run on Prolific

2

Helpfulness

N = 50

Rate model-selected
good / medium / bad
messages for each
map.



What people say · think · do

4 preregistered experiments · 1,200 explanations · 24 maps · run on Prolific

3

Baseline

N = 50

Navigate every map with no explanation as map-difficulty baseline.



What people say · think · do

4 preregistered experiments · 1,200 explanations · 24 maps · run on Prolific

4

With explanation

N = 150

Navigate with bad / medium / good messages (within-subject).

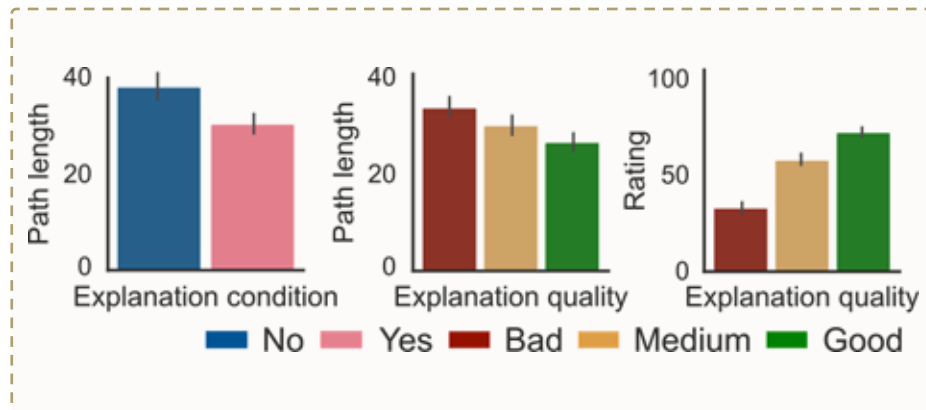


The model's ranking predicts judgments and behavior

JUDGED MORE HELPFUL

Good > Medium > Bad

a clean monotonic ordering.



The model's ranking predicts judgments and behavior

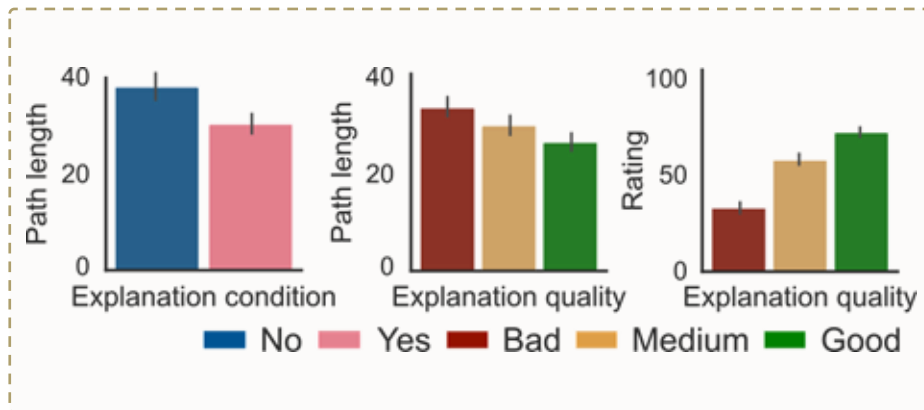
NAVIGATED MORE EFFICIENTLY

Any explanation beats none

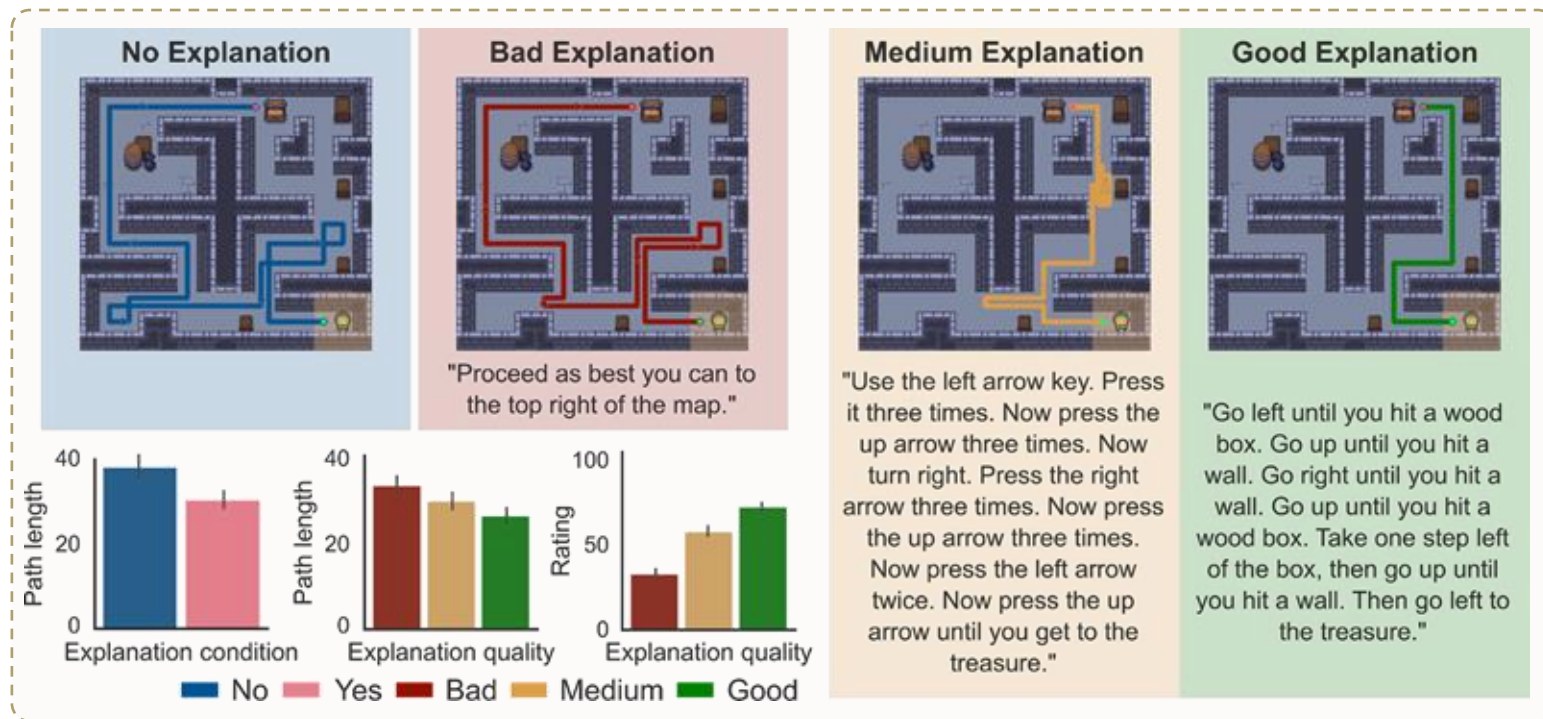
JUDGED MORE HELPFUL

Good > Medium > Bad

a clean monotonic ordering.

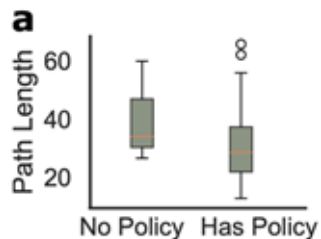


The model's ranking predicts judgments and behavior



Policy for speed, value for recovery

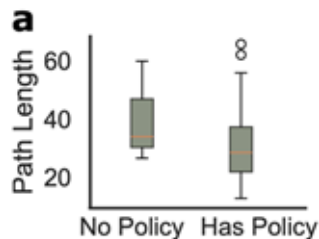
- **POLICY** for speed



Path length drops when a policy is given -> faster navigation.

Policy for speed, value for recovery

• POLICY for speed



Path length drops when a policy is given -> faster navigation.

how policy breaks

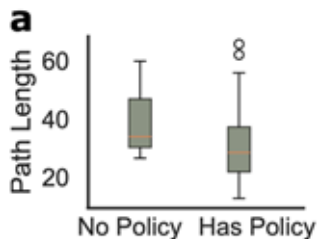
Overcomplicated

long conditional instructions exceed the compiler.

THE RESULTS

Policy for speed, value for recovery

• POLICY for speed



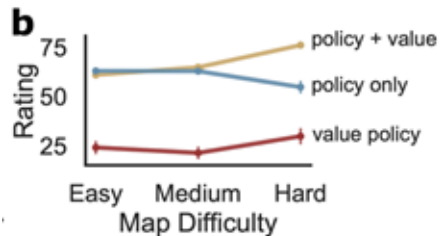
Path length drops when a policy is given -> faster navigation.

how policy breaks

Overcomplicated

long conditional instructions exceed the compiler.

• VALUE for recovery

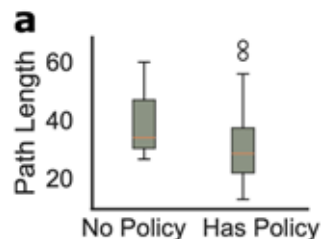


On hard maps, policy + value holds up while policy-only collapses.

THE RESULTS

Policy for speed, value for recovery

• POLICY for speed



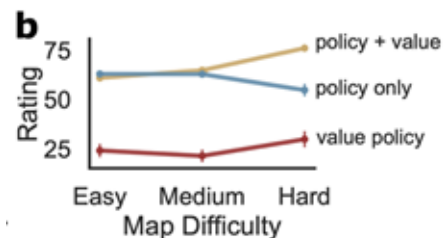
Path length drops when a policy is given -> faster navigation.

how policy breaks

Overcomplicated

long conditional instructions exceed the compiler.

• VALUE for recovery



On hard maps, policy + value holds up while policy-only collapses.

how value breaks

Spatial ambiguity a region with no referable landmark; the agent keeps re-querying.

Richer worlds and varied agents

1. Scaling map complexity

from grid dungeons to open, structured worlds (trees, water, animals) with multiple viable routes.

2. Varying the agent

new capabilities (Walk · Locate · Clamber · Vault · Hop · Forage) and a changing field-of-view.

How do explainers adapt to what a listener **KNOWS** and what they **CAN DO**?





Takeaways & Limitations

Judge an explanation by what the listener can **do** with it - a model that **acts** on explanations recovers both how people rate them and how well they navigate.

Takeaways & Limitations

Judge an explanation by what the listener can **do** with it - a model that **acts** on explanations recovers both how people rate them and how well they navigate.

1 The message carries more

Caution & effort cues: “careful, this map’s tricky”

Error signals: “if you hit the barrel, you’ve gone too far”.

Takeaways & Limitations

Judge an explanation by what the listener can **do** with it - a model that **acts** on explanations recovers both how people rate them and how well they navigate.

1 The message carries more

Caution & effort cues: “careful, this map’s tricky”

Error signals: “if you hit the barrel, you’ve gone too far”.

2 We assume perfect grounding

Our agent maps words to the world with **almost no ambiguity**.

Reference resolution: “the barrel,” “the room on the left” but which one?

Thank you

Questions are very welcome.

Hanqi Zhou

hanqi.zhou@uni-tuebingen.de



 GitHub: [cicl-stanford/
explaining_how_cogsci26](https://github.com/cicl-stanford/explaining_how_cogsci26)