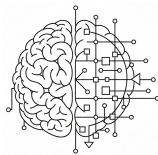


What is reasoning?



Jacqueline Maasch

Digital Life Initiative | Cornell Tech | 24 February 2026

Based on work in:

Position: Beyond *Reasoning Zombies* — AI Reasoning Requires Process Validity

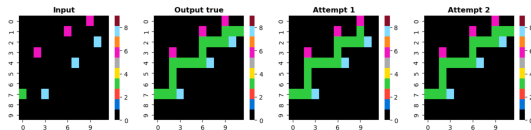
Rachel Lawrence^{1*} Jacqueline Maasch^{2*}

^{*}Equal contribution ¹Microsoft Research, Cambridge, UK ²Cornell Tech, Department of Computer Science, New York, NY. Correspondence to: Rachel Lawrence <rachel.lawrence@microsoft.com>, Jacqueline Maasch <maasch@cs.cornell.edu>.

Problem setting.

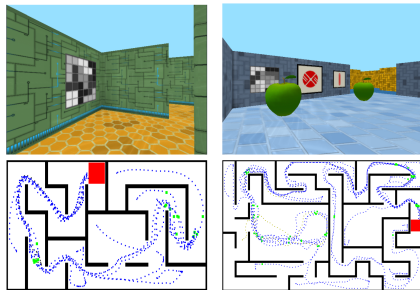
Toward reasoning machines

1 Problem setting



Learning rules OOD [ARC-AGI-1]

question string - lengths 	answer string - lengths 
Natalia sold clips to 48 of her friends in April, and then she sold half as many clips in May. How many clips did Natalia sell...	Natalia sold $48/2 = <<48/2>24>>24$ clips in May. Natalia sold $48+24 = <<48+24>72>>72$ clips altogether in April and May. #### 72
Weng earns \$12 an hour for babysitting. Yesterday, she just did 50 minutes of babysitting. How much did she earn?	Weng earns $12/60 = \$<<12/60>0.2>>0.2$ per minute. Working 50 minutes, she earned $8.2 \times 50 = \$<<8.2 \times 50>10>>10$. #### 10
Betty is saving money for a new wallet which costs \$100. Betty has only half of the money she needs. Her parents decided to give her...	In the beginning, Betty has only $100 / 2 = \$<<100/2>50>>50$. Betty's grandparents gave her $15 \times 2 = \$<<15 \times 2>30>>30$. This means...



Navigating [1]

Math question answering (QA) [GSM8K]

Toward reasoning machines

1 Problem setting

```
theorem and_swap (p q : Prop) : p ∧ q → q ∧ p := by
  intro h          -- assume p ∧ q with proof h, the goal is q ∧ p
  apply And.intro   -- the goal is split into two subgoals, one is q and the other is p
  · exact h.right    -- the first subgoal is exactly the right part of h : p ∧ q
  · exact h.left     -- the second subgoal is exactly the left part of h : p ∧ q
```

A simple proof in Lean, a language for automated theorem proving [\[wiki\]](#).

- Historical contributions from logic, formal methods, symbolic AI.
- Paradigm shift: recent progress from large reasoning models (LRMs) [\[2\]](#).

Toward reasoning machines

1 Problem setting

- Proliferation of open questions.
 - Is autonomous reasoning an emergent behavior that arises with scale [3, 4]?
 - Can LRMs be formally characterized as autonomous reasoners?
 - What are the **mechanisms** by which reasoning occurs in deep probabilistic models?
 - What continuity does LRM “reasoning” share with historical treatments of reasoning?
 - **Answers are contingent on how reasoning is defined.**

Toward reasoning machines

1 Problem setting

So then, *what is reasoning?*

And how can we link the latent construct to measurable proxies?

Operationalization. An operational definition is “a description of something in terms of the *operations* (procedures, actions, or processes) by which it could be observed and measured. For example, the operational definition of anxiety could be in terms of a test score, withdrawal from a situation, or activation of the sympathetic nervous system. The process of creating an operational definition is known as *operationalization*” [APA].

Toward operational definitions

1 Problem setting

- **Lack of consensus.**
 - Claims of emergent reasoning in generative AI are commonplace.
 - Yet, “**there is not a clear definition of what it entails**” [2].
 - Shifting goalposts arise.
 - Strong benchmark performance is often conflated with reasoning.
 - Absent operational definitions, construct validity of reasoning evaluation is **unfalsifiable**.
- **Reasoning versus superficial emulation.**
 - The blackbox design and natural language interface of LRMs present a nontrivial challenge: **differentiating true reasoning from reasoning-like speech** [5, 6, 7, 8].
 - Emulation can manifest as talking like a reasoner, with no guarantees that conclusions arose from reasoning rather than memorization, guessing, Clever Hans effects, etc [9, 10, 11].

Toward operational definitions

1 Problem setting

What's the problem?

- P1** Reasoning in generative AI has experienced unnecessary and addressable definitional ambiguity, where imprecise and overloaded definitions are often misaligned with historical treatments of this topic (when definitions are provided at all).
- P2** This breeds mismeasurement, promotes an illusion of shared understanding among researchers, and subverts measurable progress toward trustworthy AI reasoning.

Toward operational definitions

1 Problem setting

- We do not see a justification for *reinventing reasoning* in the context of generative AI.
- We project that **operational definitions that are method agnostic** (simultaneously compatible with symbolic, neural, and neuro-symbolic methods) will provide:
 - Conceptual unification for the AI reasoning community.
 - Greater research value in the long run.

Problem significance.

Why should we care?

2 Problem significance

R1 Reasoning is a necessary (but not sufficient) precondition for AGI [12].

R2 Questionable construct validity [13, 14]: are we measuring what we claim we are?

R3 Usership outpaces evidence of trustworthiness.

Why should we care?

2 Problem significance

R1 Reasoning is a necessary (but not sufficient) precondition for AGI [12].

R2 Questionable construct validity [13, 14]: are we measuring what we claim we are?

R3 Usership outpaces evidence of trustworthiness.

Process $\neq$ product

3 Construct validity & benchmarking

François Chollet [15] on the risks of “confusing the process of intelligence” (reasoning, in our case) “with the artifact produced by this process” (e.g., QA responses), ignoring the generating mechanism:

“

In the case of AI, the focus on achieving task-specific performance while placing no conditions on how the system arrives at this performance has led to systems that, despite performing the target tasks well, largely do not feature the sort of human intelligence that the field of AI set out to build.

Process $\neq$ product

3 Construct validity & benchmarking

“

*A theory of bounded rationality, then, will be as much concerned with procedural rationality, the quality of the **processes of decision**, as with substantive rationality, the **quality of the outcome**. To understand the former, one must have a theory of the psychology of the decision maker; to understand the latter, one needs have only a theory of the goal (the utility function) and the external environment. [...] When rationality is associated with reasoning **processes**, and not just with its **products**, limits on the abilities of Homo sapiens to reason cannot be ignored.*

— Herbert Simon [16]

The limitations of benchmarking

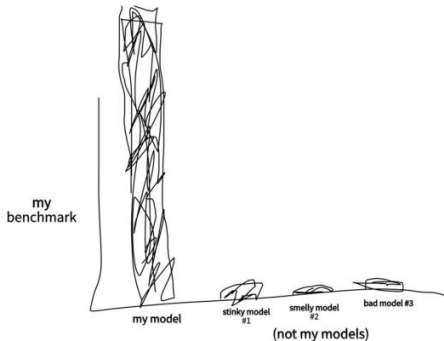
3 Construct validity & benchmarking



goosewin

@Goosewin

guys you're never gonna believe this



8:01 PM · February 14, 2026

120 Retweets 2.9K Likes

The limitations of benchmarking

3 Construct validity & benchmarking

“

Testing shows the presence, not the absence of bugs.

— Edsger Dijkstra [17]

- **From proofs to benchmarks** [13].
 - Shift toward exploratory research and empirical evaluations.
 - Shift away from hypothesis-driven confirmatory research, guarantees, formal verification.
- Benchmarks often suffer from weak construct validity [14, 18], data contamination [19], overfitting, minimal quality control, gaming, SOTA hacking, and selective reporting [20].
- AI evaluation has yet to “**mature into a proper ‘science’**” [21]

Operationalizing valid & sound reasoning.

Reasoning is a (learnable) rule-based process

4 Operationalizing valid & sound reasoning

Definition 1 (Reasoning, informal). The process of selecting and applying sequences of rules that act on prior beliefs and current evidence to obtain principled belief updates in evolving states.

Definition 2 (Reasoner, informal). A goal-oriented decision-maker that implements reasoning.

Core ingredients — intuition

4 Operationalizing valid & sound reasoning

1. Process.
2. Goal.
3. Rules.
4. Evidence.
5. Prior beliefs.
6. Current beliefs.
7. Evolving states.

Core ingredients — intuition

4 Operationalizing valid & sound reasoning

1. Process.

- Reasoning is a dynamic process with $T \geq 1$ hops, stages, time steps, or reasoning steps.
- This process implies a **design component**: sequences of rules or actions are chosen by the reasoner according to some justification.
- The process of selection is where **agency, intelligence, or creativity** may come into play, while the process of execution necessitates **exactness and rigor**.
- Note that it may be perfectly reasonable for the selection criterion to be random selection.
- Broome: reasoning is a *process*, “something [one] does,” a “rule-governed operation” [22].

2. Goal.

3. Rules.

4. Evidence.

5. Prior beliefs.

6. Current beliefs.

7. Evolving states.

Core ingredients — intuition

4 Operationalizing valid & sound reasoning

1. Process.

2. Goal.

- The reasoner executes a reasoning process to achieve some **outcome of interest**.
- Examples: the answer to a question, the solution to a puzzle, the shortest path through a maze, a mathematical proof, the optimal action to take under resource constraints, etc.
- Reasoning validity is not necessarily tied to successful attainment of a goal.
- In practice, we can encode the goal in a stopping rule.

3. Rules.

4. Evidence.

5. Prior beliefs.

6. Current beliefs.

7. Evolving states.

Core ingredients — intuition

4 Operationalizing valid & sound reasoning

1. Process.

2. Goal.

3. Rules.

- The rule set unambiguously maps the reasoning state at t to the state at $t + 1$.
- In general, **rules are learnable and defeasible**.
- Rules are selected with justification prior to deployment.
- Rules can be algorithms, formulae, theorems, axioms, laws, policies, premises, assumptions, decision boundaries, etc.
- Rules can be extrinsically imposed on the reasoner or they can be learned autonomously from data on-the-fly. Rules can be fixed or continuously updated in light of new information.

4. Evidence.

5. Prior beliefs.

6. Current beliefs.

7. Evolving states.

Core ingredients — intuition

4 Operationalizing valid & sound reasoning

1. Process.

2. Goal.

3. Rules.

4. Evidence.

- **Exogenous or extrinsically obtained information.**

- A continuous stream of data that is updated at each step t or at intervals (or, trivially provided at $t = 0$ and never updated).

- Current evidence denotes information presented at t , along with the historical record: aggregated information up to $k \geq 0$ steps prior to t .

- Evidence may be gained directly through sequential interactions with an uncertain environment (as in RL) or provided without direct collection (e.g., retrospective data).

5. Prior beliefs.

6. Current beliefs.

7. Evolving states.

Core ingredients — intuition

4 Operationalizing valid & sound reasoning

1. Process.

2. Goal.

3. Rules.

4. Evidence.

5. **Prior beliefs.**

- **Endogenous or intrinsically obtained information.**

- Prior beliefs are the outputs of previous reasoning steps (**intermediate conclusions** along the reasoning pathway that led to step t).

- Can be defeasible: beliefs can be overwritten if proven false (e.g., in backtracking proof search), refined if insufficient, or maintained and aggregated with current beliefs at step t .

- They can also be provided at $t = 0$ (e.g., initializing Bayesian priors).

6. Current beliefs.

7. Evolving states.

Core ingredients — intuition

4 Operationalizing valid & sound reasoning

1. Process.

2. Goal.

3. Rules.

4. Evidence.

5. Prior beliefs.

6. **Current beliefs.**

- **Conclusions drawn in the transition from $t - 1$ to t .**

- When $t = T$, equivalent to the **terminal conclusion** of the reasoning process.

- The nature of the terminal conclusion is a defining property of domain-specific reasoning, e.g.: the output of a function in mathematical reasoning, an optimal action in practical reasoning, a moral verdict in moral reasoning, a judiciary decision in legal reasoning, etc.

7. Evolving states.

Core ingredients — intuition

4 Operationalizing valid & sound reasoning

1. Process.
2. Goal.
3. Rules.
4. Evidence.
5. Prior beliefs.
6. Current beliefs.

7. **Evolving states.**

- A reasoner will generally maintain an **internal representation of its world state**, which updates over time.
- The existence of an external environment is also implied by our choice to model evidence as a stream of extrinsic signals.
- A well-defined concept of **external environment is not relevant in all cases** (e.g., in some mathematical reasoning domains).

Reasoning is a (learnable) rule-based process

4 Operationalizing valid & sound reasoning

Definition 3 (Reasoning, formal). Let $S_t := \langle B_t, E_t, R_t \rangle$ denote the reasoner's state at time step t , where B_t denotes current belief, E_t denotes aggregated evidence up to time t , and R_t denotes the current set of established rules. Then, *reasoning* is the iterated application over steps t of rules $r \in R_{t-1}$ to prior beliefs B_{t-1} and current evidence E_t , by which we obtain dynamically updated states S_t , and where every output B_t for $t > 0$ is the result of a rule application $r(B_{t-1}, E_t)$ to the contents of state S_{t-1} .

Reasoning is a (learnable) rule-based process

4 Operationalizing valid & sound reasoning

Reasoning components

$t \in [0, \dots, T]$	Reasoning step.
$\{\mathcal{B}_i\}_{i=0}^T, \mathcal{B}_i \in \mathbf{B}$	Beliefs.
$\{\mathcal{E}_i\}_{i=0}^T, \mathcal{E}_i \in \mathbf{E}$	Evidence.
$\{\mathcal{R}_i\}_{i=0}^T, \mathcal{R}_i \in \mathbf{R}$	Rule set.
$\mathcal{S}_i := \langle \mathcal{B}_i, \mathcal{E}_i, \mathcal{R}_i \rangle, \mathcal{S}_i \in \mathbf{S}$	States.
$\mathcal{R}_t^L := \{r \in \mathcal{R}_t \mid r : \mathbf{B} \times \mathbf{E} \rightarrow \mathbf{B}\}$	
$\mathcal{R}_t^M := \{r \in \mathcal{R}_t \mid r : \mathbf{R} \times \mathbf{B} \times \mathbf{E} \rightarrow \mathbf{R}\}$	

Reasoner components

$s_L : \mathbf{R} \times \mathbf{B} \times \mathbf{E} \rightarrow \mathcal{R}^L$	Local rule selector.
$s_M : \mathbf{R} \times \mathbf{B} \times \mathbf{E} \rightarrow \mathcal{R}^M$	Meta rule selector.
$s_{\text{stop}} : \mathbf{S} \rightarrow \{0, 1\}$	Stopping rule.
$\text{tr} : \mathbf{S} \times \mathcal{R}^L \times \mathcal{R}^M \times \mathbf{S} \rightarrow \Sigma^*$	Trace writer.
$\mathcal{T} := \{\text{tr}(\mathcal{S}_{i-1}, r_i^L, r_i^M, \mathcal{S}_i)\}_{i=1}^T$	Reasoning trace.
where $r_i^L := s_L(\mathcal{R}_t, \mathcal{B}_t, \mathcal{E}_{t+1})$ and $r_i^M := s_M(\mathcal{R}_t, \mathcal{B}_t, \mathcal{E}_{t+1})$.	

Validity & Soundness

4 Operationalizing valid & sound reasoning

Definition 4 (Validity). A transition from state S_t to S_{t+1} is *valid* if and only if it arises from the application of a rule $r \in R_t$ to components of state S_t .

Definition 5 (Soundness). A valid transition from state S_t to S_{t+1} is *sound* if and only if all premises (as encoded by beliefs, rules, and evidence) are true with respect to external evaluation.

Validity is independent of rule selection

4 Operationalizing valid & sound reasoning

“

*Correct reasoning is not reasoning you are **required** to do by rationality, but reasoning you are **permitted** to do by rationality.*

— John Broome [22]

Claim. Because validity is independent of soundness, and any properly-typed rule application creates a valid output, the validity of a reasoning process is independent of the algorithm used to select the rule sequence, regardless of external ground truth.

Valid reasoning as exact rule application

4 Operationalizing valid & sound reasoning

Algorithm 1 Valid reasoning as exact rule application.

Input. Initial rules $\mathcal{R}_0$, beliefs $\mathcal{B}_0$, evidence stream $\{\mathcal{E}_i\}_{i=1}^T$.

$\mathcal{R}, \mathcal{B}, \mathcal{E} \leftarrow \mathcal{R}_0, \mathcal{B}_0, \mathcal{E}_0$

$\mathcal{S} \leftarrow (\mathcal{R}, \mathcal{B}, \mathcal{E})$

$t \leftarrow 0$

while not $s_{\text{stop}}(\mathcal{S})$ **do**

$\mathcal{E}' \leftarrow \mathcal{E}_{t+1}$

$r^L \leftarrow s_L(\mathcal{R}, \mathcal{B}, \mathcal{E}') \text{ \{Select local rule.\}}$

$\mathcal{B}' \leftarrow r^L(\mathcal{B}, \mathcal{E}') \text{ \{Apply local rule, update beliefs.\}}$

$r^M \leftarrow s_M(\mathcal{R}, \mathcal{B}', \mathcal{E}') \text{ \{Select meta rule.\}}$

$\mathcal{R}' \leftarrow r^M(\mathcal{R}, \mathcal{B}', \mathcal{E}') \text{ \{Apply meta rule, update rules.\}}$

$\mathcal{S}' \leftarrow (\mathcal{R}', \mathcal{B}', \mathcal{E}')$

$\mathcal{T}.\text{append}(\text{tr}(\mathcal{S}, r^L, r^M, \mathcal{S}')) \text{ \{Update trace.\}}$

$\mathcal{R}, \mathcal{B}, \mathcal{E}, \mathcal{S} \leftarrow \mathcal{R}', \mathcal{B}', \mathcal{E}', \mathcal{S}'$

$t += 1$

end while

Return $\mathcal{B}, \mathcal{T}$

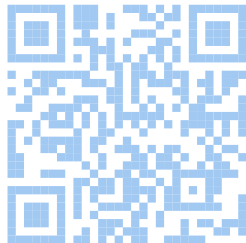
Mapping to domain-specific cases

4 Operationalizing valid & sound reasoning

Example	Beliefs $\mathcal{B}_t$	Evidence $\mathcal{E}_t$	Rules $\mathcal{R}_t$	Goal / Conclusion
Logical deduction	Derived formulas in the proof state Γ .	Premises E_0 are given at $t = 0$ and not updated thereafter.	Fixed inference rules (e.g., modus ponens, introduction/elimination rules)	Derive a target formula φ such that $\Gamma \vdash \varphi$.
Bayesian inference	Current posterior $p(\theta \mid D_{1:t})$	Newly observed data D_t (possibly aggregated with prior observations).	Bayes' rule and any auxiliary update rules (e.g., conjugate prior updates, approximation schemes).	Obtain updated posterior beliefs $p(\theta \mid D_{1:t})$.
Reinforcement learning	Current value function estimates, policy parameters, and internal state representations.	Observed environment states, state transitions, and rewards obtained by interaction with the environment.	Update rules such as temporal difference or policy gradient updates, plus any meta rules adapting learning rates or architectures.	Learn a policy that maximizes expected return (i.e., select approximately optimal actions over time).
Nonmonotonic logic	Current set of accepted conclusions, including defeasible ones.	New information that may conflict with existing conclusions (e.g., exceptions, defaults).	Nonmonotonic inference rules that support belief revision and retraction, plus meta rules for revising the rule set itself.	Maintain a coherent, defeasible belief set that updates appropriately under new, possibly contradictory evidence.
Turing machine	Current configuration of the machine: tape contents $\sigma_t \in \Sigma^*$, head position h_t , and control state $q_t \in Q$, collectively encoded as $\mathcal{B}_t = \langle \sigma_t, h_t, q_t \rangle$.	Input word $w \in \Sigma^*$ (typically fixed at $t = 0$).	Transition relation or function $\delta_t : Q \times \Sigma \rightarrow Q \times \Sigma \times \{L, R\}$.	Compute the value of a (partial) function $f : \Sigma^* \rightarrow \Sigma^*$ on input w , i.e., reach a halting configuration with output tape σ_T such that σ_T encodes $f(w)$.
Probabilistic next-token prediction	Current token given prior tokens.	Token(s) provided at initialization (e.g., the first word of a sentence, the prompt to an instruction-tuned LLM, etc.).	Probability rules (e.g., chain rule), structural assumptions (e.g., Markovianity), maximum likelihood formulae.	Iterate procedure until query is answered (e.g., string is of desired length, etc.).



View the project page:



`maasch@cs.cornell.edu` ★ `https://jmaasch.github.io/`

References

- [1] P. Mirowski et al. “Learning to Navigate in Complex Environments”. In: *International Conference on Learning Representations*. 2017.
- [2] J. Huang et al. “Towards Reasoning in Large Language Models: A Survey”. In: *Findings of the Association for Computational Linguistics: ACL 2023*. 2023, pp. 1049–1065.
- [3] J. Wei et al. “Emergent Abilities of Large Language Models”. In: *Transactions on Machine Learning Research* (2022).
- [4] J. González et al. “Does reasoning emerge? examining the probabilities of causation in large language models”. In: *Advances in Neural Information Processing Systems (NeurIPS 2024)* (2024).
- [5] P. Shojaee et al. “The illusion of thinking: Understanding the strengths and limitations of reasoning models via the lens of problem complexity”. In: *arXiv preprint arXiv:2506.06941* (2025).
- [6] A. Hüyük et al. “Reasoning Elicitation in Language Models via Counterfactual Feedback”. In: *International Conference on Learning Representations* (2025).
- [7] J. Maasch et al. “Compositional causal reasoning evaluation in language models”. In: *International Conference on Machine Learning* (2025).
- [8] X. Xu et al. “RE-IMAGINE: Symbolic Benchmark Synthesis for Reasoning Evaluation”. In: *International Conference on Machine Learning*. 2025.

References

- [9] S. Lapuschkin et al. “Unmasking Clever Hans predictors and assessing what machines really learn”. In: *Nature communications* 10.1 (2019), p. 1096.
- [10] J. Kauffmann et al. “Explainable AI reveals Clever Hans effects in unsupervised learning models”. In: *Nature Machine Intelligence* (2025), pp. 1–11.
- [11] M. Mitchell. “Artificial intelligence learns to reason”. In: *Science* 387.6740 (2025), eadw5211.
- [12] D. Hendrycks et al. “A Definition of AGI”. In: *arXiv preprint arXiv:2510.18212* (2025).
- [13] M. Herrmann et al. “Position: Why We Must Rethink Empirical Research in Machine Learning”. In: *International Conference on Machine Learning*. PMLR. 2024, pp. 18228–18247.
- [14] H. Wallach et al. “Position: Evaluating Generative AI Systems Is a Social Science Measurement Challenge”. In: *Forty-second International Conference on Machine Learning Position Paper Track*. 2025.
- [15] F. Chollet. “On the Measure of Intelligence”. In: *arXiv preprint arXiv:1911.01547* (2019).
- [16] H. A. Simon. “Bounded rationality in social science: Today and tomorrow”. In: *Mind & Society* 1.1 (2000), pp. 25–39.
- [17] E. W. Dijkstra. *Software Engineering Techniques*. Ed. by J. Buxton et al. 1970.
- [18] A. Alaa et al. “Position: Medical Large Language Model Benchmarks Should Prioritize Construct Validity”. In: *Forty-second International Conference on Machine Learning Position Paper Track*. 2025.
- [19] C. White et al. “LiveBench: A Challenging, Contamination-Limited LLM Benchmark”. In: *The Thirteenth International Conference on Learning Representations*. 2025.

References

- [20] Z. Cheng et al. “Benchmarking is Broken—Don’t Let AI be its Own Judge”. In: *Advances in Neural Information Processing Systems (NeurIPS 2025)* (2025).
- [21] L. Weidinger et al. “Toward an evaluation science for generative ai systems”. In: *arXiv preprint arXiv:2503.05336* (2025).
- [22] J. Broome. *Rationality Through Reasoning*. John Wiley & Sons, 2013.
- [23] H. S. Richardson. “Moral Reasoning”. In: *The Stanford Encyclopedia of Philosophy*. Ed. by E. N. Zalta. Fall 2018. Metaphysics Research Lab, Stanford University, 2018.
- [24] N. Kolodny. “Why be rational?” In: *Mind* 114.455 (2005), pp. 509–563.
- [25] J. Broome. “The unity of reasoning?” In: *Spheres of reason* 1.9 (2009), pp. 62–93.
- [26] P. Hieronymi. “The use of reasons in thought (and the use of earmarks in arguments)”. In: *Ethics* 124.1 (2013), pp. 114–127.
- [27] F. Portoraro. “Automated Reasoning”. In: *The Stanford Encyclopedia of Philosophy*. Ed. by E. N. Zalta et al. Summer 2025. Metaphysics Research Lab, Stanford University, 2025.
- [28] G. Wang et al. “Hierarchical Reasoning Model”. In: *arXiv preprint arXiv:2506.21734* (2025).

Supplemental slides.

Alternative definitions

Reasoning, in Stanford Encyclopedia of Philosophy [23]. Active or explicit thinking, in which the reasoner, responsibly guided by her assessments of her reasons [24] and of any applicable requirements of rationality [25, 22], attempts to reach a well-supported answer to a well-defined question [26].

Automated reasoning, in Stanford Encyclopedia of Philosophy [27]. Reasoning is the ability to make inferences [by] proving the conclusion from the given assumptions by the systematic application of **rules** of deduction embedded within the reasoning program.

Alternative definitions

Reasoning [28]. The process of devising and executing complex goal-oriented action sequences.

Reasoning [22]. Reasoning is a mental process in which you operate on the contents of your attitudes, following a **rule**.

Reasoning [2]. Reasoning is the process of thinking about something in a logical and systematic way, using evidence and past experiences to reach a conclusion or make a decision.