

Analogy making as amortised model construction

David G. Nagy^{1,2}, Tingke Shen², Hanqi Zhou^{1,2}, Charley M. Wu^{1,2,3,4,†},
Peter Dayan^{2,1,†}

david.nagy@uni-tuebingen.de

¹University of Tübingen, Germany

²Max Planck Institute for Biological Cybernetics, Germany

³Technical University of Darmstadt, Darmstadt, Germany

⁴Hessian.AI, Darmstadt, Germany

† Equal contribution

Abstract

Humans flexibly construct internal models to navigate novel situations. To be useful, these internal models must be sufficiently faithful to the environment that resource-limited planning leads to adequate outcomes; equally, they must be tractable to construct in the first place. We argue that analogy plays a central role in these processes, enabling agents to reuse solution-relevant structure from past experiences and amortise the computational costs of both model construction (construal) and planning. Formalising analogies as partial homomorphisms between Markov decision processes, we sketch a framework in which abstract modules, derived from previous construals, serve as composable building blocks for new ones. This modular reuse allows for flexible adaptation of policies and representations across domains with shared structural essence.

1 Introduction

Humans, like artificial reinforcement learning (RL) agents, maintain internal representations of their environment, enabling them to interpret sensory inputs and predict action outcomes. In RL, these representations usually take the form of Markov decision processes (MDPs¹) — abstract formalisations of the environment that support planning and decision making. Typically, formulating such an MDP — defining the states, observations, actions, rewards — is performed manually by an external designer (although recent approaches allow the learning of this model from experience in restricted domains; [Schrittwieser et al., 2020](#); [Hafner et al., 2025](#)). For example, a vacuum-cleaning robot might perceive its environment in terms of ‘obstacles,’ ‘dirt concentration,’ and ‘returning to the charging dock’—rather than ‘minimal yet cozy living rooms,’ ‘ringing phones,’ or ‘following one’s dreams’—reflecting the designer’s assumptions about which abstractions are relevant for vacuuming floors.

Unlike a robot vacuum confined to navigating an apartment, humans must operate in vastly more diverse situations: navigating not just an apartment, but an entire city (on foot, by bike, or by public transport); decorating an apartment or building a shelter, working in an office, negotiating, interpreting social nuances, or solving math problems. Crucially, the choice of the abstract representation used in these situations is deeply intertwined with decision-making processes operating over it — an effective representation can render difficult problems trivially easy to solve, while a poor one can make a solution impossible ([Giunchiglia & Walsh, 1992](#); [Abel, 2022](#); [Ravindran & Barto, 2003](#)).

¹We use MDP here as a shorthand for the broader family of Markov decision process-based models, including extensions such as POMDPs (partially observable MDPs), BAMDPs (Bayes-adaptive MDPs), and IPOMDPs (interactive POMDPs), which allow for uncertainty, learning, and multi-agent reasoning, respectively.

We argue that the vast diversity of situations humans encounter rules out reliance on a pre-designed MDP — or even a predetermined set of MDPs. Instead, humans must be capable of constructing MDPs for novel situations *on-demand*. In other words, they must fill the shoes of not just the user, but also the designer, of their own internal MDPs — yielding what we refer to throughout this paper as situation specific *construals*, following Ho et al. (2022). Both these roles have to be played in the face of limited cognitive resources. Unfortunately, the obvious solution of simplification leads to variants of the notoriously difficult “frame problem” (Dennett, 1990; Icard & Goodman, 2015).

Here, we take inspiration from a long body of work focusing on the central role of analogy and metaphor in human cognition (Hofstadter & Sander, 2013; Lakoff & Johnson, 2008; Gentner & Hoyos, 2017). We propose that analogies (understood broadly) underlie the human ability for on-demand model construction, and more generally offer a powerful mechanism by which RL agents can flexibly adapt past knowledge to novel circumstances. Specifically, we formalise potential analogies as mappings between MDPs that preserve solution-relevant structure (and in some cases the entire solution). Through such analogical mappings, humans can construct new internal models by adapting and combining abstract MDPs that were effective in past situations. This reuse enables the transfer of prior computations, effectively amortising both the construction and solution costs of models for novel situations (Gershman & Goodman, 2014; Dasgupta & Gershman, 2021).

As building blocks of analogies, we propose that the brain extracts structural regularities across diverse situations in the form of reusable fragments of MDPs, and incrementally compiles them into a library of consistently useful *modules*. As fragments are adapted and reused across increasingly varied contexts, they become progressively abstracted, gaining broader applicability as sources of analogy. Over time, such modules may become decoupled from their original contexts entirely, forming widely reusable conceptual primitives such as ‘door’, ‘stairs’, ‘fire’ or ‘clock’ (Fig. 1).

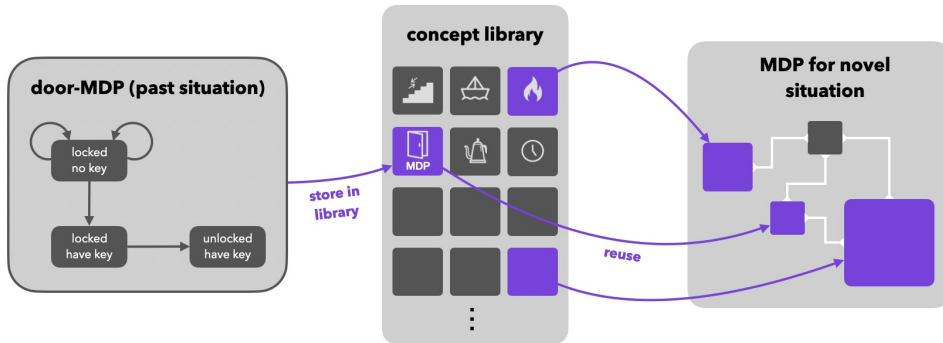


Figure 1: Library of abstract MDP modules for amortised model construal.

Consider a small child observing that the front door (previously thought part of the wall) can be unlocked and opened by inserting and turning a specific object — the *key*. This transforms it into a functional door, similar to those inside the house. Crucially, only a specific key will work; other objects, even if similarly shaped, will fail. Years later, the child is given their first email account. A parent explains, through an explicit analogy, that the password is the ‘key’ for their account — an ‘*email-key*’. Of course, the analogy is imperfect, e.g. there is no ‘turning’ the password, only clicking on the login button. Yet, the analogy likely helps the child greatly. For instance, they now understand that no other password, even of similar length or characters is likely to open the account. They know to guard the password, because if someone steals it they will gain access to their messages. But if they do want to grant someone access, they can simply share the password.

Despite surface differences in terms of low-level actions and states, the key/door and the email/password situations share high-level structure. By seeing the login screen as a ‘door’ and using the password as a ‘key’, the child can reuse a familiar mental representation, including transition dynamics, possible actions, and useful policies. The analogy preserves the situation’s essence, allowing the child to transfer knowledge across domains.

In the following, we formulate the challenge of on-demand model construction, proposing analogy as a mechanism for amortising both the construction and solution costs of MDPs. We argue that humans construct models not only by analogy to single past situations, but also by recomposing abstract modules drawn from an internal library that they concurrently grow. We formalise analogies as partial MDP homomorphisms that preserve structure relevant for planning and decision making. We end with an overview of key computational challenges inherent in forming, maintaining and using such a library of composable modules.

2 The challenge of on-demand model construction

Unlike typical artificial RL agents, humans must not only use, but also construct their internal models of the environment. To build our formalisation of this joint challenge, we contrast the vacuum-cleaning robot with the child from our earlier examples.

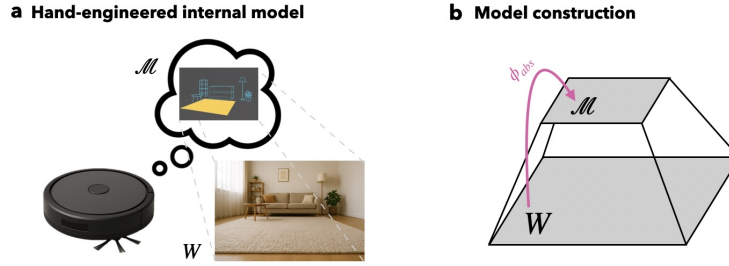


Figure 2: The design process of certain real world artefacts (e.g. vacuum-cleaning robots) resembles solving the computational problem of MDP construal.

Consider first the perspective of a robot tasked with efficiently cleaning an apartment: it must cover all accessible floor areas, avoid bumping into furniture, and return to its charging station before the battery is depleted. In an RL setting, the subset of environmental knowledge maintained by the robot to perform this task is encoded in what we call its internal MDP $\mathcal{M}$, created by its designer. This internal MDP defines the states s_t (e.g., the robot’s position, or whether a room is clean), actions a_t (e.g., move forward, return to charger), the transition function $T(s_t, a_t, s_{t+\Delta t})$ (e.g., walls are impassable, the dock recharges the battery), and the reward function R (e.g., for successfully cleaning a room, which could depend on the selected cleaning mode such as ‘quick clean’ or ‘deep clean’). Concretely, the robot must compute and execute a policy $\pi_{\mathcal{M}}^*$:

$$\pi_{\mathcal{M}}^*(a_t|s_t) = \arg \max_{\pi} \mathbb{E}[R|\pi, \mathcal{M}] \quad (1)$$

Now let us shift perspective to the engineer responsible for designing the robot’s internal representation of the environment $\mathcal{M}$ in the first place. The first challenge for the designer is to ensure that the policy $\pi_{\mathcal{M}}^*$ that the robot derives using its internal MDP survives contact with the real world. Thus, the designer must keep in mind that the optimisation in Eq. 1 is only a proxy for the true objective of maximising reward *in the real world*, represented by a hypothetical ground-truth MDP W :

$$\pi_W^* = \arg \max_{\pi} \mathbb{E}[R'|\pi, W] \quad (2)$$

Ideally, the internal model $\mathcal{M}$ would perfectly match the world, i.e., $\mathcal{M} = W$. However, as Jorge Luis Borges’ famous parable illustrates, a map as detailed as the territory it represents is of little practical use (Borges, 1960). Likewise, an internal model that mirrors every detail of the environment would be computationally intractable. Therefore the second challenge for the engineer is to ensure that operating the model is tractable under the robot’s limited computational resources. Formally², the cost of computation C_s required to arrive at a solution has to stay below the maximum

²Here, we use ‘solution’ in a broad sense (e.g., it may refer to computing a plan via model-based methods, learning value estimates or policies via model-free RL, or hybrid approaches that combine both), while treating computational cost in abstract terms, which may include time and energy required for solving the model.

computational capacity $C_{\max}(\text{robot})$. Taken together, we can formalise the engineer’s objective as:

$$\mathcal{M}^* = \arg \max_{\mathcal{M}} \mathbb{E}[R' | \pi_{\mathcal{M}}^*, W], \quad s.t. \ C_s(\mathcal{M}) \leq C_{\max}(\text{robot}), \quad (3)$$

where R' represents the real-world goals that the engineer aims to achieve. A key difficulty is that there is no a priori bound on what aspects of the real world need to be included in the model to solve a given problem. In principle, anything could be related to anything and it is difficult to judge which details can be safely ignored. This is exactly a version of the frame problem (Dennett, 1990; Icard & Goodman, 2015), which has a long and storied history in philosophy and AI, and cognitive science more broadly (Butz et al., 2025).

In summary, the engineer’s objective is to construct an internal MDP for the robot that is tractable to solve, and whose resulting policy leads to the desired outcomes when executed in the real world. Related ideas have been explored in the context of human planning, such as in the framework of *value-guided construals* (Ho et al., 2022), and in the context of the frame problem (Icard & Goodman, 2015). In the following, we refer to the resulting $\mathcal{M}^*$ as the *construal*, although it is sometimes also referred to as a submodel (Icard & Goodman, 2015), situation model (Zwaan & Radvansky, 1998), or mental space (Fauconnier, 1994). The process of building the construal is sometimes referred to as model synthesis (Wong et al., 2023a; Ahmed et al., 2025).

However, an essential challenge is often ignored, namely that optimising Eq. 3 itself incurs computational costs, denoted C_c . For the vacuum-cleaning robot, we might intuitively think of C_c as the resources expended by the engineer in designing the model, including the engineer-hours available for the design. This additional constraint is especially important to consider if we now take the perspective of the human child, whom we described as both the designer and user of their own mental models. Unlike the robot, where the roles of model construction and model use are split, the human must manage both processes internally. Thus, we can define a combined objective for model construction we refer to as the *on-demand objective*:

$$\mathcal{M}^* = \arg \max_{\mathcal{M}} \mathbb{E}[R'' | \pi_{\mathcal{M}}^*, W], \quad s.t. \ C_s(\mathcal{M}) + C_c(\mathcal{M}, W) \leq C_{\max}(\text{child}). \quad (4)$$

That is, to act effectively in novel situations, the child must construct an internal model $\mathcal{M}$ that is not only useful but also computationally feasible to both create and to solve.

3 Analogy making as amortised model construction

How do humans contend with the challenge of on-demand model construction? We argue that a key strategy is to reuse previously computed construals, by drawing analogies between past and present situations. These analogies amortise previous computations (Dasgupta & Gershman, 2021; Gershman & Goodman, 2014; Huys et al., 2015), enabling new environments to be modelled in old terms, and previously computed successful policies to be mapped, at least approximately, to new situations. Central to this process is a library of abstract, composable fragments (i.e., modules) from which new construals can be flexibly assembled (Fig. 1; Zhou et al., 2024; Rubino et al., 2023). The contents of the library can be seen as offering a form of powerful inductive bias for the construal process, as well as including a complete, or at least a partial, solution of the resulting construal. In turn, the library is populated and refined by means of refactoring and organizing the construals as they are created. The computational cost of maintaining this library can be viewed as an upfront investment that reduces the effort required to construct future models.

In the following, we sketch an account of how analogical mappings can support efficient model construction. First, building on the example of the child learning to use their email account, we examine a simplified case in which the construal for a new situation is built from a single abstract module. Then, we extend our discussion to scenarios where multiple modules must be composed and their respective solutions integrated into a coherent whole. Finally, we consider how such modules might be extracted, refined, abstracted, and maintained over time.

3.1 Making a single analogy

Analogy as partial MDP homomorphism. Recall our earlier example, in which the parent’s guidance helped a child draw an analogy from a familiar ‘door–key’ scenario (the source) to the new ‘email–password’ situation (the target). We formalise such analogies as structure-preserving mappings ϕ between a source $\mathcal{M}_{\text{source}}$ and a target construal $\mathcal{M}_{\text{target}}$. More specifically, we consider an analogy to be a *partial (S)MDP homomorphism* (Ravindran & Barto, 2002; 2003), consisting of a state mapping $f(s)$ and an action mapping $g_s(a)$. The analogy might map the locked state of a door to the logged-out state of an email account ($f(s_{\text{locked, no key}}) = \tilde{s}_{\text{logged out, no pwd}}$) and the act of picking up the key to recalling a password ($g_s(a_{\text{get key}}) = \tilde{a}_{\text{recall pwd}}$; Fig. 3a). The mapping qualifies as a strict homomorphism if it preserves the reward and transition structures of the source MDP. Intuitively, this can be seen as a generalisation of the deterministic homomorphism conditions to the stochastic case, where the homomorphism must commute with the system dynamics: applying the dynamics and then the homomorphism yields the same result as applying the homomorphism first and then the abstract dynamics (Hartmanis, 1966). Crucially, analogies typically involve partial mappings (Gentner, 1983; Gentner & Hoyos, 2017; Lakoff & Johnson, 2008), preserving only aspects of structure deemed relevant in the current context. We refer to $\phi = (f(s), g_s(a))$ as a *partial MDP homomorphism* when the conditions are only required to hold over a relevant subset of the state-action space rather than globally over the entire MDP.

Construal by analogy. From the perspective of on-demand model construction, making an analogy involves the child constructing a new internal model $\mathcal{M}_{\text{target}}$ by using $\mathcal{M}_{\text{source}}$ as an inductive bias and the real-world task W_{target} as a constraint (Fig. 3b). For simplicity, we assume that the child already knows W_{target} (i.e., has access to the relevant low-level dynamics) but lacks effective abstractions to cope with its complexity. In more realistic scenarios where the low-level MDP must also be learned, the abstract structure provided by analogies could additionally guide exploration. Constructing the new MDP can be seen as attempting to reuse the high-level structure of the source $\mathcal{M}_{\text{source}}$ for the target construal of W_{target} , but with redefined state and action abstraction functions. Thus, in this simplified setting when the analogy source is explicitly provided, the central computational challenges are to work out the entirety of ϕ_{analogy} mapping the abstract models, and the new ϕ'_{abs} mapping low level target states and actions to the abstract model’s macro-states and macro-actions (Fig. 3c). When a parent says, ‘the password is like a key’, we hypothesise that they are effectively supplying part of these mappings, implying for example that ‘unlocking’ has to do with typing the password on the keyboard. This partial mapping lowers the computational burden for the child, making it easier to infer the rest (Fig. 3c). In cases without external instruction, the child might find new analogies by searching over their pre-existing library of abstract modules and considering candidates according to the above process, which may also be amortised (Ellis et al., 2023; O’Donnell et al., 2009).

Mapping solutions through analogies. A key benefit of construal through analogy is that it enables amortisation of the planning process (Huys et al., 2015), notably when action abstractions from the source module are also effective for planning in the new construal. More generally, we propose that modules in the library encode not just MDP fragments, but also associated partial solutions (e.g., policies or policy fragments) that can also be transferred via analogy. For example, interpreting a login screen as a locked door immediately constrains a child’s policy search: actions like clicking ‘login’ without a password are expected to be futile, and many unproductive strategies can be ruled out. Instead, the child’s search can be biased toward solutions involving password acquisition, providing an inductive constraint that narrows the solution space even when the exact details are different.

Importantly, generalisation to new situations often arises through much more mundane analogies. Interpreting the door of a previously unknown building—or even a novel object like a spaceship—as an instance of the abstract ‘door’ module immediately suggests a range of reasonable actions: to unlock, open, close, lock, go through, change the lock, or even bang on it, but not, for example, to walk through it while it is closed, attempt to unlock it with an arbitrary key, or try to plead with it. Drawing an analogy through a module to represent a new object (e.g., identifying it as a new

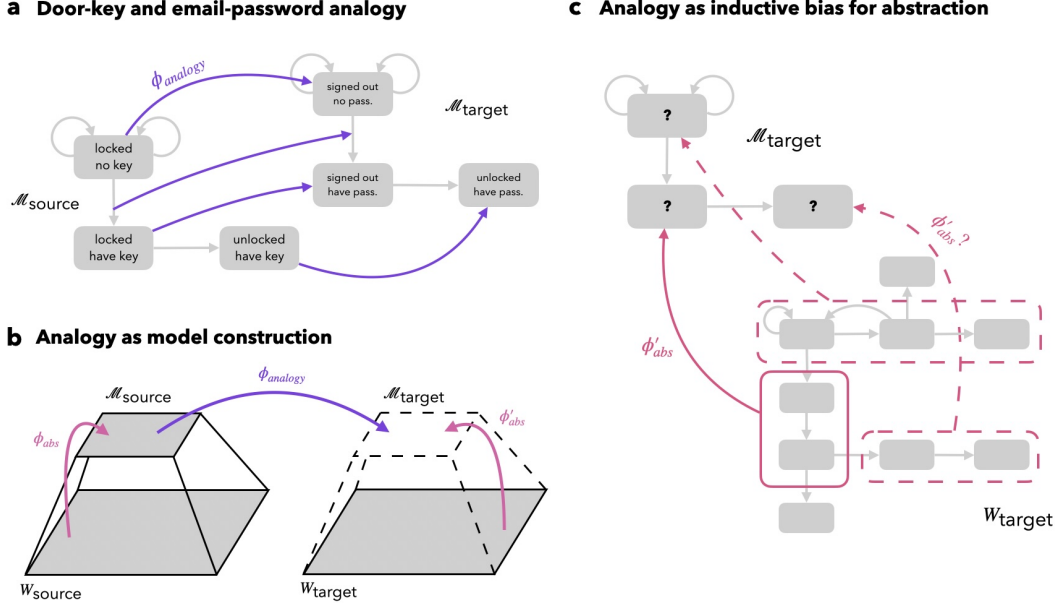


Figure 3: Analogy-guided model construction. a) The child draws an analogy $\phi_{analogy}$ between the construals for the door-key ($\mathcal{M}_{source}$) and email-password ($\mathcal{M}_{target}$) situations. b) In our framework, this means constructing $\mathcal{M}_{target}$ using $\mathcal{M}_{source}$ as an inductive bias, constrained by the real-world task W_{target} . Here, the parent explicitly points out a previous MDP as the source. More generally the child has to search over their library of modules (see Fig. 4). c) In making the analogy, the child has to work out which micro-states and micro-actions correspond to which abstract states and actions. Parental guidance provides partial state-action mappings (solid pink arrows), helping the child work out the remaining ones (dashed pink arrows).

instance of a 'door') is closely related to perceiving its *affordances* (Gibson, 2014). In Gibson's sense, the environment is perceived in terms of the possibilities for action it offers to the agent. An affordance, like an action abstraction associated with a module, reflects not only properties of the object (or environment), but also the capabilities of the agent: for example, a lake is 'support' that affords standing-on for a water bug but not for a human, and a tall bench might be a 'chair' that affords sitting for an adult but not for a small child. From this perspective, the cognitive mechanism that maps parts of the environment to MDP modules can also be understood as a mechanism for affordance perception.

Ravindran & Barto (2002) show that, in straightforward cases where part of the abstract structure is shared exactly, if $\pi_{source}(s, a)$ is an optimal policy for the source MDP, then the mapped policy $\pi_{target}(f(s), g_s(a))$ is also optimal for the target. Losses in performance arising from only approximate homomorphisms can be quantified via Bounded-parameter MDPs (Ravindran & Barto, 2002).

3.2 Model construction through composition of modules

Thus far, we have considered model construction via analogy with a single module. However, the diversity of real-world situations demands a more flexible strategy for reuse: the ability to decompose previously formed construals into fragments, and to recombine these fragments in novel configurations. By compiling a library of MDP modules $\mathcal{L} = \{\mu_i\}$ that preserve solution-relevant structure, it is possible to acquire an abstract vocabulary for model construction as well as the solution of the resulting models. Following previous work viewing amortisation via a library as an inference problem (O'Donnell et al., 2009; Ellis et al., 2023; Zhao et al., 2023; Bowers et al., 2023; Zhou et al., 2024;

Rule et al., 2020), we decompose the derivation of the policies $\pi_{1:t}^*$, the construals $\mathcal{M}_{1:t}$ and the library $\mathcal{L}$ on the basis of $W_{1:t}$ into three separate stages that the agent alternates between (Fig. S1):

1. $\mathcal{M}_t | \mathcal{L}, W_t$ — Conditional on a fixed library $\mathcal{L}$, what is the construal $\mathcal{M}_t$ for W_t ?
2. $\pi_t^* | \mathcal{M}_t, \mathcal{L}$ — Conditional on a fixed construal and the library, what is an effective policy?
3. $\mathcal{L} | \mathcal{M}_{1:t}, \pi_{1:t}^*$ — Conditional on the history of construals and solutions, what modules should be in the library?

In the following, we briefly discuss key computational challenges raised by the above questions.

Composing modules and adapting solutions. Since analogies can map between arbitrary subcomponents of an MDP, construals for novel situations may be assembled piecewise, drawing components from distinct previous situations ($\mathcal{M}_{t_i}$) or abstract modules (μ_i). For example, a child familiar with their apartment may interpret an office door by analogy with their apartment door, while drawing on their experience with a school projector to understand a similar device in a conference room. This compositional reuse allows useful state and action abstractions to be imported in parts, enabling efficient planning in unfamiliar situations, echoing recent proposals that modular building blocks support zero-shot generalisation in novel environments (Bakermans et al., 2025).

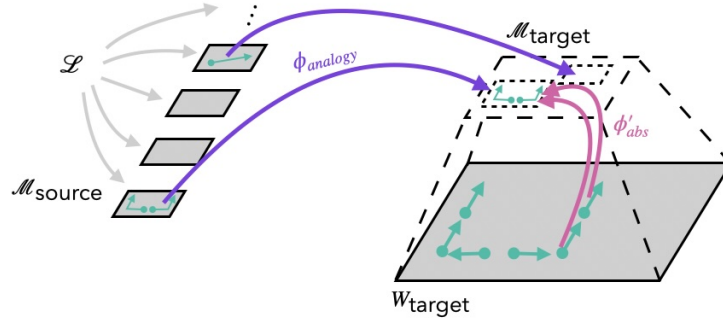


Figure 4: Modular construction of model, and transfer of solutions through the analogy.

A key challenge in composing modules is that while each module encodes an approximate solution to the frame problem within its original context, there is no guarantee that these solutions will transfer smoothly to new configurations. For instance, while a chair might consistently afford sitting across all of a child’s previous experiences, in a novel context—perched on a pedestal in a museum—the same actions may lead to undesirable outcomes. Therefore, the transferability of policies and action abstractions associated with a given module (e.g., the chair), or even the choice of the module from the library, can also be dependent on other modules in the construal (e.g., the pedestal). These relations between modules may be possible to represent and learn separately, as in prior approaches to compositional MDPs (Diuk et al., 2008; Guestrin et al., 2003; Bapst et al., 2019).

While we technically distinguish between the stages of construal and solution, these processes are tightly intertwined. A good construal can render the solution trivial, while a poor one may make a solution intractable or even impossible. For example, in a new situation, the optimal policy might rely entirely on a single transferred action abstraction (which involves mapping the action abstraction into the target domain’s low level actions), making policy derivation straightforward and effectively reducing the problem to creating the construal. Conversely, if the imported abstractions are more granular or less applicable to the new situation, more of the computational burden is deferred to the solution phase. In either case, creating the construal is not merely a precursor to solving the problem, but often the first, and most critical step in the solution process itself.

Module extraction and refinement. While partial MDP homomorphisms allow reuse of components directly from past construals stored in memory, this process is made more efficient by extracting consistently useful fragments into abstract, stand-alone modules. By keeping track of the

aspects of the module that are frequently discarded in mapping to new situations, one might refine the module to serve as better source of analogies. For example, while the 'key' module might originally be extracted into a module that applies to various physical key-like objects (car key, room key, etc.), in the course of forming more distant analogies (including with passwords), the child will likely also form a highly abstract 'key' module that allows them to interpret sentences such as 'the key insight in the theory is...'. Such a movement of the source of an analogy over time from a richly represented, concrete situation to one of a set of highly abstract concepts has been aptly described as the 'career of a metaphor' by [Bowdle & Gentner \(2005\)](#). Formally, the increasing abstraction of amortised components resembles the evolution of a newly defined primitive in fragment grammars ([O'Donnell et al., 2009](#)) and in approaches to program induction ([Bowers et al., 2023](#)). In a neuroscience context, it might be understood as the consolidation of memory traces from episodic to semantic memory ([Káli & Dayan, 2004](#); [Lengyel & Dayan, 2007](#); [Nagy et al., 2025](#)).

4 Discussion

We argued that on-demand model construction poses a uniquely demanding computational challenge. While humans appear to solve this challenge with remarkable flexibility, our account suggests that this flexibility is grounded not in general-purpose reasoning, but in powerful inductive biases derived from past experience, encoded as reusable modules and structured analogies. The severity of the on-demand model construction challenge—compounded by the frame problem—likely necessitates amortisation of computational costs not only within individuals, but across individuals, societies, and generations ([Wu et al., 2022](#)). We propose that analogies enable especially useful and abstract building blocks of construals and policies to become embedded in cultural products—stories, idioms, images, and physical artifacts ([Vélez et al., 2022](#)). This allows the creation of a cultural repository of reusable abstractions, and consequently ways of seeing situations and solving problems ([Wu et al., 2024](#)).

Despite recent advances in learning world models for RL agents ([Schrittwieser et al., 2020](#); [Hafner et al., 2025](#)), current methods still struggle with flexible adaptation to novel situations. Approaches focused narrowly on task-relevant representations often fail to generalise even across minor reconstructions of the same environment, though reconstruction or self-supervised objectives can offer some improvement ([Anand et al., 2021](#); [Hafner et al., 2025](#)). However, even modest changes in goals, object identity or layout can still cause severe drops in performance. In contrast, theory-based RL (TBRL) methods can sometimes support the kind of structured generalisation observed in humans ([Tsividis et al., 2021](#); [Pouncy & Gershman, 2022](#); [Pouncy et al., 2021](#)), but typically aim to recover a relatively complete model of the environment's dynamics rather than simplified, task-specific construals — thereby side-stepping the frame problem, but scaling poorly in more complex domains. More recent TBRL methods have begun to construct abstract simulators in limited domains by offloading parts of the abstraction process to large language models (e.g. [Ahmed et al., 2025](#); [Wong et al., 2023b;a](#)).

While we have proposed a conceptual framework for analogy-guided model construction, much remains to be done to translate this sketch into a concrete computational theory. Key challenges include formalizing the processes by which modules are extracted, composed, and abstracted over time, and identifying mechanisms for searching over analogies that are both flexible and computationally tractable. In particular, it remains an open question how agents might balance the reuse of familiar structure against the discovery of genuinely novel construals in environments where no clear analogy applies. Nevertheless, if modular reuse through analogy underlies the human ability to construct useful models under resource constraints, then formalising these mechanisms may offer a path toward more robust artificial agents.

Acknowledgments

The authors thank Mihály Bányai for helpful discussions and comments on the manuscript, and Ryutaro Uchiyama, Tianyuan Teng, Rahul Bhui, Lionel Wong, Tyler Brooke-Wilson and Joshua

Tenenbaum for insightful discussions. DGN, HZ, & CMW are supported by the German Federal Ministry of Education and Research (BMBF): Tübingen AI Center, FKZ: 01IS18039A, funded by the Deutsche Forschungsgemeinschaft (DFG, German Research Foundation) under Germany’s Excellence Strategy–EXC2064/1–390727645, and funded by the DFG under Germany’s Excellence Strategy – EXC 2117 – 422037984.

References

- David Abel. A theory of abstraction in reinforcement learning. *CoRR*, 2022.
- Zergham Ahmed, Joshua B. Tenenbaum, Christopher J. Bates, and Samuel J. Gershman. Synthesizing world models for bilevel planning, March 2025. arXiv:2503.20124 [cs].
- Ankesh Anand, Jacob Walker, Yazhe Li, Eszter Vértés, Julian Schrittwieser, Sherjil Ozair, Théophane Weber, and Jessica B. Hamrick. Procedural Generalization by Planning with Self-Supervised World Models, November 2021. arXiv:2111.01587 [cs].
- Jacob J. W. Bakermans, Joseph Warren, James C. R. Whittington, and Timothy E. J. Behrens. Constructing future behavior in the hippocampal formation through composition and replay. *Nature Neuroscience*, pp. 1–12, March 2025. ISSN 1546-1726. Publisher: Nature Publishing Group.
- Victor Bapst, Alvaro Sanchez-Gonzalez, Carl Doersch, Kimberly Stachenfeld, Pushmeet Kohli, Peter Battaglia, and Jessica Hamrick. Structured agents for physical construction. In *Proceedings of the 36th International Conference on Machine Learning*, pp. 464–474. PMLR, May 2019. ISSN: 2640-3498.
- Jorge Luis Borges. Del rigor en la ciencia. *El hacedor*, 45, 1960.
- Brian F Bowdle and Dedre Gentner. The career of metaphor. *Psychological review*, 112(1):193, 2005.
- Matthew Bowers, Theo X. Olausson, Lionel Wong, Gabriel Grand, Joshua B. Tenenbaum, Kevin Ellis, and Armando Solar-Lezama. Top-Down Synthesis for Library Learning. *Proceedings of the ACM on Programming Languages*, 7(POPL):1182–1213, January 2023. ISSN 2475-1421.
- Martin V Butz, Maximilian Mittenbühler, Sarah Schwöbel, Asya Achimova, Christian Gumbsch, Sebastian Otte, and Stefan Kiebel. Contextualizing predictive minds. *Neuroscience & Biobehavioral Reviews*, 168:105948, 2025.
- Ishita Dasgupta and Samuel J. Gershman. Memory as a Computational Resource. *Trends in Cognitive Sciences*, 25(3):240–251, March 2021. ISSN 1364-6613.
- D Dennett. Cognitive Wheels: the Frame Problem of AI. *The philosophy of artificial intelligence*, 1990.
- Carlos Diuk, Andre Cohen, and Michael L. Littman. An object-oriented representation for efficient reinforcement learning. In *Proceedings of the 25th international conference on Machine learning - ICML ’08*, pp. 240–247, Helsinki, Finland, 2008. ACM Press. ISBN 978-1-60558-205-4.
- Kevin Ellis, Lionel Wong, Maxwell Nye, Mathias Sablé-Meyer, Luc Cary, Lore Anaya Pozo, Luke Hewitt, Armando Solar-Lezama, and Joshua B. Tenenbaum. DreamCoder: growing generalizable, interpretable knowledge with wake–sleep Bayesian program learning. *Philosophical Transactions of the Royal Society A: Mathematical, Physical and Engineering Sciences*, 381(2251):20220050, July 2023. ISSN 1364-503X, 1471-2962.
- Gilles Fauconnier. *Mental spaces: Aspects of meaning construction in natural language*. Cambridge University Press, 1994.
- Dedre Gentner. Structure-mapping: A theoretical framework for analogy. *Cognitive Science*, 7(2): 155–170, April 1983. ISSN 0364-0213.

- Dedre Gentner and Christian Hoyos. Analogy and Abstraction. *Topics in Cognitive Science*, 9(3):672–693, 2017. ISSN 1756-8765. [_eprint: https://onlinelibrary.wiley.com/doi/pdf/10.1111/tops.12278](https://onlinelibrary.wiley.com/doi/pdf/10.1111/tops.12278).
- S. Gershman and Noah D. Goodman. Amortized Inference in Probabilistic Reasoning. *Cognitive Science*, 2014.
- James J. Gibson. The Theory of Affordances: (1979). In *The People, Place, and Space Reader*. Routledge, 2014. ISBN 978-1-315-81685-2. Num Pages: 5.
- Fausto Giunchiglia and Toby Walsh. A theory of abstraction. *Artificial Intelligence*, 57(2):323–389, October 1992. ISSN 0004-3702.
- Carlos Guestrin, Daphne Koller, Chris Gearhart, and Neal Kanodia. Generalizing plans to new environments in relational MDPs. In *Proceedings of the 18th international joint conference on Artificial intelligence, IJCAI’03*, pp. 1003–1010, San Francisco, CA, USA, August 2003. Morgan Kaufmann Publishers Inc.
- Danijar Hafner, Jurgis Pasukonis, Jimmy Ba, and Timothy Lillicrap. Mastering diverse control tasks through world models. *Nature*, pp. 1–7, 2025.
- Juris Hartmanis. *Algebraic structure theory of sequential machines (prentice-hall international series in applied mathematics)*. Prentice-Hall, Inc., 1966.
- Mark K. Ho, David Abel, Carlos G. Correa, Michael L. Littman, Jonathan D. Cohen, and Thomas L. Griffiths. People construct simplified mental representations to plan. *Nature*, 606(7912):129–136, June 2022. ISSN 1476-4687. Number: 7912 Publisher: Nature Publishing Group.
- Douglas R. Hofstadter and Emmanuel Sander. *Surfaces and Essences: Analogy as the Fuel and Fire of Thinking*. Basic Books, April 2013. ISBN 978-0-465-02158-1. Google-Books-ID: XkQT5eTnurYC.
- Quentin JM Huys, Níall Lally, Paul Faulkner, Neir Eshel, Erich Seifritz, Samuel J Gershman, Peter Dayan, and Jonathan P Roiser. Interplay of approximate planning strategies. *Proceedings of the National Academy of Sciences*, 112(10):3098–3103, 2015.
- Thomas F Icard and Noah Goodman. A Resource-Rational Approach to the Causal Frame Problem. *Proceedings of the 37th Annual Meeting of the Cognitive Science Society*, pp. 6, 2015.
- Szabolcs Káli and Peter Dayan. Off-line replay maintains declarative memories in a model of hippocampal-neocortical interactions. *Nature neuroscience*, 7(3):286–294, 2004.
- George Lakoff and Mark Johnson. *Metaphors We Live By*. University of Chicago Press, December 2008. ISBN 978-0-226-47099-3. Google-Books-ID: r6nOYYtxzUoC.
- Máté Lengyel and Peter Dayan. Hippocampal contributions to control: the third way. *Advances in neural information processing systems*, 20, 2007.
- David G. Nagy, Gergő Orbán, and Charley M. Wu. Adaptive compression as a unifying framework for episodic and semantic memory. *Nature Reviews Psychology*, pp. 1–15, June 2025. ISSN 2731-0574.
- Timothy J O’Donnell, Noah D Goodman, and Joshua B Tenenbaum. Fragment Grammars: Exploring Computation and Reuse in Language. 2009.
- Thomas Pouncy and Samuel J. Gershman. Inductive biases in theory-based reinforcement learning. *Cognitive Psychology*, 138:101509, November 2022. ISSN 00100285.
- Thomas Pouncy, Pedro Tsividis, and Samuel J. Gershman. What Is the Model in Model-Based Planning? *Cognitive Science*, 45(1), January 2021. ISSN 0364-0213, 1551-6709. DOI: 10.1111/cogs.12928. URL <https://onlinelibrary.wiley.com/doi/10.1111/cogs.12928>.

- Balaraman Ravindran and Andrew G. Barto. Model Minimization in Hierarchical Reinforcement Learning. In Sven Koenig and Robert C. Holte (eds.), *Abstraction, Reformulation, and Approximation*, pp. 196–211, Berlin, Heidelberg, 2002. Springer. ISBN 978-3-540-45622-3.
- Balaraman Ravindran and Andrew G. Barto. SMDP homomorphisms: an algebraic approach to abstraction in semi-Markov decision processes. In *Proceedings of the 18th international joint conference on Artificial intelligence, IJCAI’03*, pp. 1011–1016, San Francisco, CA, USA, August 2003. Morgan Kaufmann Publishers Inc.
- Valerio Rubino, Mani Hamidi, Peter Dayan, and Charley M Wu. Compositionality under time pressure. In M. Goldwater, F. Anggoro, B. Hayes, and D. Ong (eds.), *Proceedings of the 45th Annual Conference of the Cognitive Science Society*, Sydney, Australia, 2023. Cognitive Science Society.
- Joshua S Rule, Joshua B Tenenbaum, and Steven T Piantadosi. The child as hacker. *Trends in cognitive sciences*, 24(11):900–915, 2020.
- Julian Schrittwieser, Ioannis Antonoglou, Thomas Hubert, Karen Simonyan, Laurent Sifre, Simon Schmitt, Arthur Guez, Edward Lockhart, Demis Hassabis, Thore Graepel, et al. Mastering atari, go, chess and shogi by planning with a learned model. *Nature*, 588(7839):604–609, 2020.
- Pedro A Tsividis, Joao Loula, Jake Burga, Nathan Foss, Andres Campero, Thomas Pouncy, Samuel J Gershman, and Joshua B Tenenbaum. Human-level reinforcement learning through theory-based modeling, exploration, and planning. *arXiv preprint arXiv:2107.12544*, 2021.
- Natalia Vélez, Charley M. Wu, and Fiery A. Cushman. Representational exchange in social learning: Blurring the lines between the ritual and instrumental. *Behavioral and Brain Sciences*, 45:e271, 2022.
- Lionel Wong, Gabriel Grand, Alexander K. Lew, Noah D. Goodman, Vikash K. Mansinghka, Jacob Andreas, and Joshua B. Tenenbaum. From Word Models to World Models: Translating from Natural Language to the Probabilistic Language of Thought, June 2023a. arXiv:2306.12672 [cs].
- Lionel Wong, Jiayuan Mao, Pratyusha Sharma, Zachary S. Siegel, Jiahai Feng, Noa Korneev, Joshua B. Tenenbaum, and Jacob Andreas. Learning adaptive planning representations with natural language guidance, December 2023b. arXiv:2312.08566.
- Charley M Wu, Natalia Vélez, Fiery A Cushman, Irene Cogliati Dezza, Eric Schulz, and Charley M Wu. Representational exchange in human social learning. *The Drive for Knowledge*, pp. 169–192, 2022.
- Charley M Wu, Rick Dale, and Robert D Hawkins. Group coordination catalyzes individual and cultural intelligence. *Open Mind*, 8:1037–1057, 2024.
- Bonan Zhao, Christopher G. Lucas, and Neil R. Bramley. A model of conceptual bootstrapping in human cognition. *Nature Human Behaviour*, pp. 1–12, October 2023. ISSN 2397-3374. Publisher: Nature Publishing Group.
- Hanqi Zhou, David G. Nagy, and Charley M. Wu. Harmonizing Program Induction with Rate-Distortion Theory, May 2024. arXiv:2405.05294 [cs, math, stat].
- Rolf A Zwaan and Gabriel A Radvansky. Situation models in language comprehension and memory. *Psychological bulletin*, 123(2):162, 1998.

Supplementary Materials

The following content was not necessarily subject to peer review.

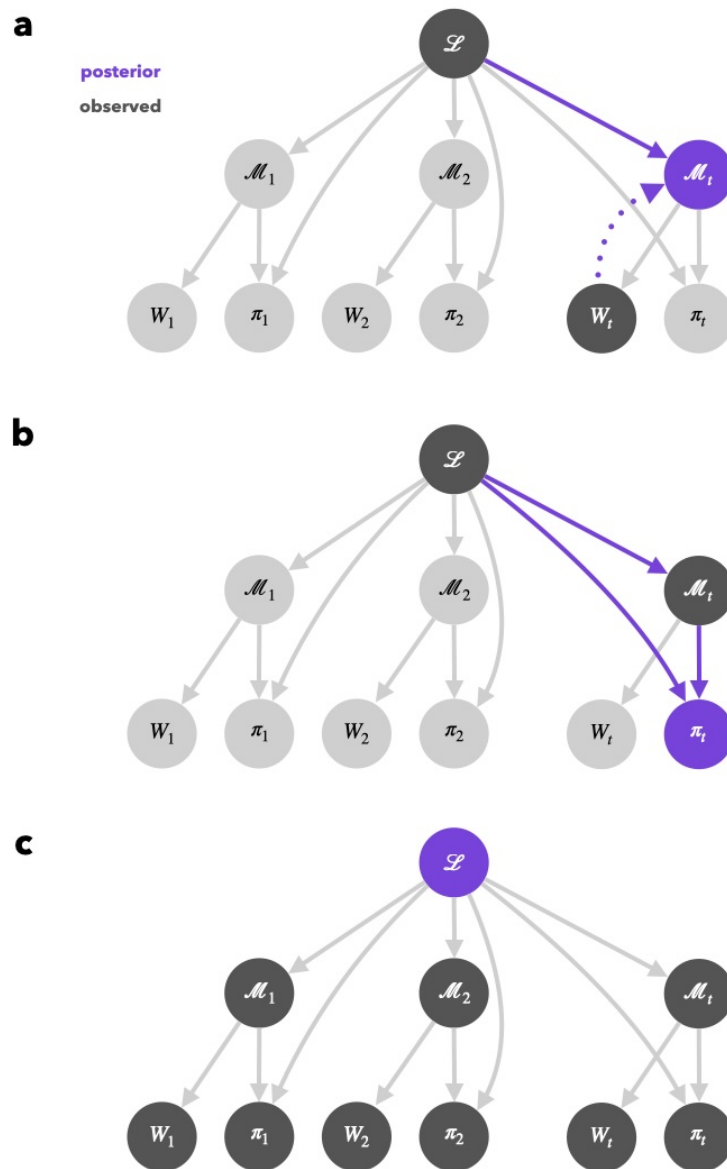


Figure S1: Library building as bayesian inference.