

Causal Process Models: Reframing Dynamic Causal Graph Discovery as a Reinforcement Learning Problem

Turan Orujlu

University of Tübingen & MPI for Biological Cybernetics

TURAN.ORUJLU@TUEBINGEN.MPG.DE

Christian Gumbsch

University of Amsterdam

Martin V. Butz

University of Tübingen

Charley M. Wu

TU Darmstadt & Hessian.AI

Editors: Bijan Mazaheri and Niels Richard Hansen

Abstract

Most neural models of causality assume static causal graphs, failing to capture the dynamic and sparse nature of physical interactions where causal relationships emerge and dissolve over time. We introduce the Causal Process Framework and its neural implementation, Causal Process Models (CPMs), for learning sparse, time-varying causal graphs from visual observations. Unlike traditional approaches that maintain dense connectivity, our model explicitly constructs causal edges only when objects actively interact, dramatically improving both interpretability and computational efficiency. We achieve this by casting dynamic interaction-graph construction for world modeling as a multi-agent reinforcement learning problem, where specialized agents sequentially decide which objects are causally connected at each timestep. Our key innovation is a structured representation that factorizes object and force vectors along three learned dimensions (mutability, causal relevance, and control relevance), enabling the automatic discovery of semantically meaningful encodings. We demonstrate that a CPM significantly outperforms dense graph baselines on physical prediction tasks, particularly for longer horizons and varying object counts.

Keywords: Causal World Models, Causal Reinforcement Learning, Causal Processes, Causal Representation Learning

1. Introduction

Causality plays a fundamental role in building intelligent systems capable of physical reasoning (Gerstenberg et al., 2020). Explicitly modeling causal relationships is increasingly recognized to be crucial for developing robust, generalizable, and interpretable neural network models capable of accurate prediction and effective intervention (Xia et al., 2021). Despite their black-box nature, models such as transformers have demonstrated surprising capacity for causal reasoning (Nichani et al., 2024; Shou et al., 2023; Melnychuk et al., 2022; Dettki et al., 2025). One explanation posits that this is possible due to the attention mechanism forming implicit causal edges between tokens (Vaswani et al., 2017; Rohekar et al., 2023). However, recent work has highlighted a phenomenon known as *over-squashing*, in which the attention mechanism (and related message-passing mechanisms in Graph Neural Networks) loses sensitivity to individual tokens or nodes (Barbero et al., 2024a; Alon and Yahav, 2021; Barbero et al., 2024b; Giovanni et al., 2023, 2024; Topping et al., 2022; Scarselli

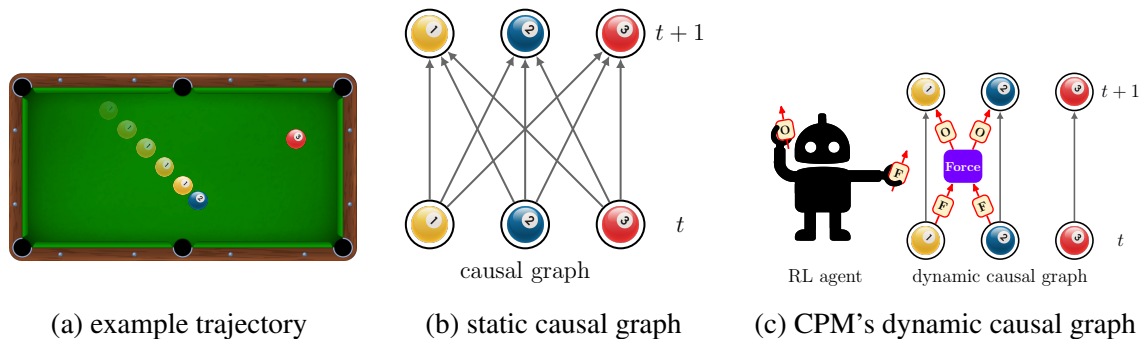


Figure 1: **Dynamic Causality**: (a) In many physical domains, such as a game of billiards, objects interact only sparsely. (b) Static causal graphs must encode *all possible* interactions, resulting in dense connectivity that fails to capture this local sparsity. (c) In a Causal Process Model (CPM), an RL agent dynamically constructs a causal graph by connecting forces and objects through process blocks, yielding a sparse, dynamic causal graph that reflects the actual interactions at each timestep.

et al., 2009; Battaglia et al., 2018). This compression of information in transformer models can sever causal chains, thus limiting the effectiveness of causal inference.

In contrast, graphical causal models, such as Pearl’s Structural Causal Models (SCMs; Pearl, 2009), explicitly encode causal relationships and thus preserve perfect causal connectivity by design. Yet, a key challenge for SCMs is *causal discovery* (Schölkopf et al., 2021): inferring the causal graph from data. Most existing approaches assume access to a complete dataset and construct a static causal graph, e.g., for all possible interactions of three billiard balls a dense graph is necessary (see Fig. 1a). This assumption is misaligned with the nature of physical environments, where causal influence is typically local in space and sparse in time (Butz, 2017; Pitis et al., 2020; Seitzer et al., 2021; Gumbsch et al., 2021; Lange and Kording, 2025). For instance, objects may only interact upon contact. Recent work has therefore emphasized the importance of *local causal models* (Pitis et al., 2020; Seitzer et al., 2021; Urpí et al., 2024; Lei et al., 2024; Willig et al., 2025) that explain causal connections through the sparsest possible graph, which changes dynamically over time.

Our work aims to bridge these areas by proposing a novel causal framework tailored to capture the dynamics of physical object interactions. We propose **Causal Process Models (CPMs)**, as a neural implementation of this framework **casting the construction of sparse dynamic causal graphs as a sequential reinforcement learning (RL) problem**. Instead of relying on dense message passing (e.g., soft attention or standard GNNs, Fig. 1b), CPMs use RL agents to dynamically determine all-or-nothing connections between entities (Fig. 1c). This allows the model to adaptively control connectivity based on the input, avoiding the over-squashing problem and enabling more efficient and interpretable causal reasoning.

Our novel causal framework is designed specifically for modeling the dynamics of physical object interactions, aiming to synthesize the formal rigor of *static dependency* theories, e.g. Pearl’s do-calculus (Pearl, 2009), with the intuitive strengths of *process-based* accounts (Russell, 1948; Salmon, 1984; Skyrms, 1981; Dowe, 2000, see Section 2 below). Our approach explicitly addresses the limitations of Pearlian SCMs by enabling the construction of sparse, time-varying causal graphs that reflect only the active interactions between objects. When modeling two colliding balls for instance, our framework only instantiates a direct causal link between the balls upon contact, for

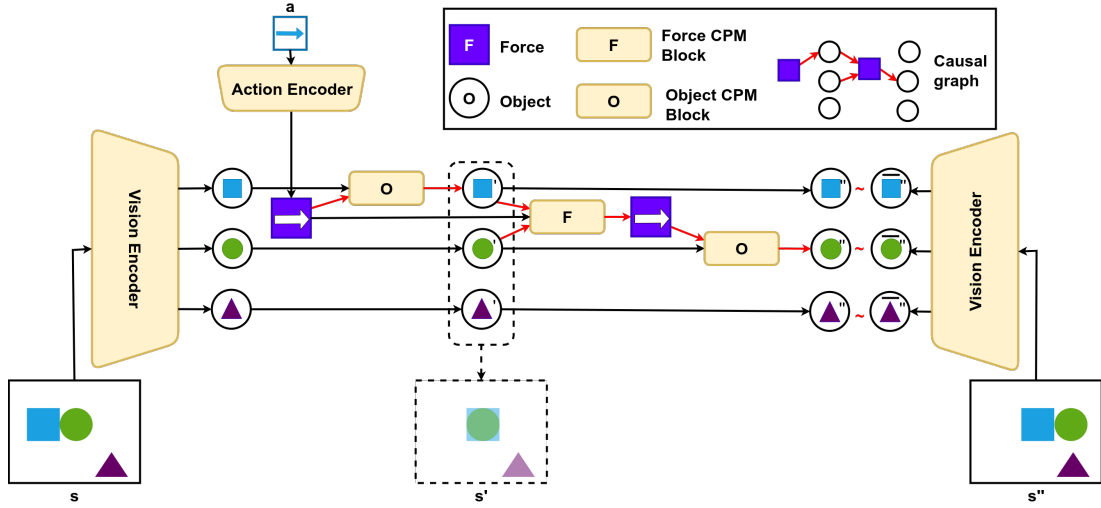


Figure 2: **Model Overview:** Our model has three components: a vision encoder, an action encoder, and a transition function. The transition function is an implementation of a *Causal Process Model*. The state is factorized into distinct object representations, actions are mapped to force representations that act as causal interventions, and the directed edges are causal. <https://github.com/torujlu/CausalProcessModels>

the transfer of momentum, while leaving them causally disconnected otherwise. This yields a computationally efficient model, only scaling with actual rather than all potential interactions, and one that is highly interpretable since the causal graph mirrors intuitive physical processes.

Our main contributions are: 1) We formalize a *Causal Process Framework* (CPF) for local causal modeling in physical environments. 2) We implement this in a neural architecture as a *Causal Process Model* (CPM) to dynamically infer sparse, time-varying causal graphs by framing edge selection as an RL problem. 3) We apply our CPM to physical interaction scenarios, demonstrating superior performance, interpretability, and scalability compared to densely connected models.

2. Related Work

2.1. Causal frameworks

Pearl’s (2009) framework of Structural Causal Models (SCMs) is a dominant approach to causal modeling, by representing causal relationships using directed acyclic graphs (DAGs). An SCM can be described as a tuple $\mathcal{C} := (\mathbf{S}, \mathbb{P}(\mathbf{U}))$ where $\mathbb{P}$ is a distribution over the exogenous variables $\mathbf{U}$ (i.e., variables external to the system and not caused by any variable within it) and $\mathbf{S}$ is a collection of structural equations of the form:

$$V_i = f_{V_i}(\mathbf{Pa}_{V_i}, \mathbf{U}_{V_i}).$$

Each endogenous variable V_i is determined by a function of its parent variables $\mathbf{Pa}_{V_i}$ (i.e., other variables in the system that directly influence V_i) and its associated exogenous noise term $\mathbf{U}_{V_i}$.

While successful in many domains, standard SCMs require extensions to handle systems characterized by dynamic object interactions; without such extensions, they fail to adequately capture the temporal and structural intricacies of such systems (Rubenstein et al., 2016; Weber, 2016; Blom et al.,

2020; Boeken and Mooij, 2024). Consider the simple scenario of two colliding balls shown in Fig. ?? . Representing this within a traditional SCM framework often requires specifying potential causal links between all properties of all objects at all relevant timescales. This leads to densely connected causal graphs (Fig. ??), with the number of causal edges scaling quadratically with time, even when interactions are sparse in reality. Such dense representations suffer from high computational costs for inference and learning, and crucially, obscure the underlying causal structure, hindering interpretability. Thus, a core challenge is to adapt standard SCMs to dynamically represent only the relevant interactions as they occur, rather than needing to specify all potential dependencies.

Recognizing these limitations, other lines of research offer valuable perspectives, often aligning closely with *causal process theories* (Russell, 1948; Salmon, 1984; Skyrms, 1981; Dowe, 2000). Research in cognitive science, such as Gerstenberg et al. (2020)’s counterfactual simulation models, leverage simulation to assess causality and responsibility in physical events, capturing process-like intuitions. Furthermore, philosophical inquiries into causal processes provide rich conceptual foundations, distinguishing *causal* processes from *pseudo*-processes by focusing on mechanisms like causal lines (Russell, 1948), defining causality in ontological terms (Salmon, 1984), or using conserved quantities (Skyrms, 1981; Dowe, 2000). However, this philosophical tradition lacks the computational formalism required for direct implementation in ML systems. Our Causal Process Framework bridges this gap by providing a computationally tractable formalism that integrates process-based intuitions with graphical causal models, enabling dynamic and sparse representations suitable for learning from visual data in physical environments.

2.2. Neural Causal Models

While philosophical causal process theories offer intuitive insights into dynamic physical interactions, their abstract nature limits direct application in scalable machine learning systems. To operationalize these ideas computationally, researchers have sought to integrate causal process intuitions with neural architectures, particularly by embedding SCMs into deep learning frameworks. Previous attempts to reconcile deep learning with SCMs have resulted in Neural Causal Models (NCMs), which model f_{V_i} as feedforward neural nets parametrized by θ_{V_i} (Xia et al., 2021). Yet this solution still suffers from the disadvantage of needing to train arbitrarily many feedforward neural networks for each node across time. To address this parameter explosion, Zecevic et al. (2021) have tried to theoretically quantify the capacity for GNNs to implement SCMs, but are restricted to the assumption of static causal graph. In contrast, Melnychuk et al. (2022) designed a Causal Transformer that incorporates temporal dynamics to infer causality over time, yet is still unable to yield interpretable graph representations. This limitation arises from its reliance on the potential outcomes framework (Rubin, 1978; Robins and Hernan, 2008), which focuses on estimating counterfactual outcomes without explicitly representing causal relationships as graphs, thus making it less suitable for discovering and utilizing sparse, time-varying structures.

2.3. Causal Reinforcement Learning

Buesing et al. (2019) have tried to take advantage of the Pearlian causality framework by reformulating the MDP graph as an SCM using which they designed a counterfactually-guided policy search. A similar approach has been pursued by Gasse et al. (2023) in which they draw parallels between confounding variables and offline RL. Neither of these approaches factors the MDP state space

into distinct object-centric nodes and their causal relations, instead focusing on the aforementioned inherent causality of the MDP structure as suggested by [Bareinboim et al. \(2021\)](#).

2.4. RL for causal discovery

Several works formulate causal structure learning as a reinforcement-learning problem, where the policy searches over graph structures to optimize a dataset-level score for a *single* global causal graph (e.g., score-based objectives over tabular variables). [Zhu et al. \(2020\)](#) and [Duong et al. \(2025\)](#) use RL to explore the combinatorial space of adjacency structures, while CORE considers an active discovery setting with interventions ([Sauter et al., 2024](#)). In contrast, our CPM uses RL *within* a predictive world model to construct sparse, time-varying interaction graphs conditioned on the current visual state. Our output is not a single global DAG but a per-timestep interaction structure used for forward prediction and downstream model-based control.

3. Causal Process Framework

Pearl’s structural causal models (SCMs) and do-calculus ([Pearl, 2009](#)) provide a powerful foundation for causal reasoning. However, without extensions, it is not straightforward to apply SCM to dynamic physical systems requiring object-centric representations and real-time causal interactions (see Appendix A for a demonstration). Prior approaches ([Buesing et al., 2019](#); [Gasse et al., 2023](#)) have attempted to bridge model-based RL and causality by representing the full Markov Decision Process (MDP) state s^t using a single node and modeling actions as direct interventions in a static causal graph. However, this approach is limited because it circumvents the problem of inferring the causal structure that generates the underlying environment dynamics (i.e., the causal context; [Butz et al., 2025](#)), and focuses only on the causal implications of action sequences.

3.1. Causal Process Models (CPMs)

To address the inability of SCMs to capture sparse, time-varying interactions, the computational burden of dense connectivity, and the loss of causal information in over-squashed message-passing, we introduce Causal Process Models (CPMs). CPMs dynamically construct sparse causal graphs that represent only active interactions, enabling both computational efficiency and interpretable causal structure in physical environments.

We adopt an *object-centric factorization* of states, in which physical objects are represented separately as object nodes $\mathcal{O} = \{O_1, O_2, \dots, O_N\}$ (e.g., balls) and interactions between objects are represented as force nodes $\mathcal{F} = \{F_1, F_2, \dots, F_M\}$ (e.g., collisions). At each timestep t , we have object states $\mathcal{O}^t = \{O_1^t, \dots, O_N^t\}$ and force states $\mathcal{F}^t = \{F_1^t, \dots, F_M^t\}$. The key insight is that not all objects interact at all times, hence we need to dynamically determine which causal edges are active.

3.1.1. DYNAMIC CAUSAL GRAPH CONSTRUCTION

To dynamically determine active causal edges in the graph, CPMs employ two types of specialized controller functions: *interaction scope controllers* $\rho_{\mathcal{O}}^t$ determine which objects interact, that is, exchange forces (e.g., based on spatial proximity), while *effect attribution controllers* $\rho_{\mathcal{O} \leftrightarrow \mathcal{F}}^t$ determine how objects are affected by these interacting forces.

Formally, each controller outputs probabilistic distributions over possible edge subsets at each timestep. The interaction scope controllers $\rho_{\mathcal{O}}^t$ define a distribution over edge sets $J^t \subseteq \mathcal{O}^t \times \mathcal{F}^t$

conditioned on current object states $\mathcal{O}^t$, yielding $J^t \sim \rho_{\mathcal{O}}^t(\cdot \mid \mathcal{O}^t)$. Similarly, the effect attribution controllers $\rho_{\mathcal{O} \leftrightarrow \mathcal{F}}^t$ define a distribution over edge sets $I^t \subseteq \mathcal{F}^t \times \mathcal{O}^{t+1}$ conditioned on current object and force states $(\mathcal{O}^t, \mathcal{F}^t)$, yielding $I^t \sim \rho_{\mathcal{O} \leftrightarrow \mathcal{F}}^t(\cdot \mid \mathcal{O}^t, \mathcal{F}^t)$.

Within this framework, state evolutions are governed by object and force update functions $f_{\mathcal{O}}$ and $f_{\mathcal{F}}$, which propagate information along the dynamically selected causal edges:

$$\begin{aligned} \text{Forces} \quad F_j^t &:= f_{\mathcal{F}} \left(F_j^{t-1}, \{O_i^{t-1}\}_{i \mid (i,j) \in J^{t-1}} \right) & \text{s.t. } J^{t-1} &\sim \rho_{\mathcal{O}}^{t-1}, \\ \text{Objects} \quad O_i^t &:= f_{\mathcal{O}} \left(O_i^{t-1}, \{F_j^t\}_{j \mid (j,i) \in I^{t-1}} \right) & \text{s.t. } I^{t-1} &\sim \rho_{\mathcal{O} \leftrightarrow \mathcal{F}}^{t-1}. \end{aligned} \quad (1)$$

Thus, update functions $f_{\mathcal{F}}$ and $f_{\mathcal{O}}$ are force- and object-specific (respectively) and invariant to the number of inputs (i.e., size of the parent node set). When interventions occur at time step $\tilde{t}$, they introduce perturbations over object nodes, denoted as $\text{do}(F_{*}^{\tilde{t}}, * \rightarrow \imath)$ representing externally applied action F_{*} on object $O_{\imath}$ (e.g., hitting a billiard ball). Intuitively, this step captures how such external influences propagate forward in time, akin to resimulating the physical system from the intervention point onward: the model dynamically recomputes the subgraph for all subsequent time steps $t \geq \tilde{t}$ by resampling causal edges and updating node states via the controllers and transition functions, ensuring the causal graph reflects the altered dynamics (following Algorithm 1 in Appendix B).

3.1.2. CONCRETE EXAMPLE: TWO COLLIDING BALLS

To illustrate, consider a scenario with three balls, where Ball 1 collides with Ball 2 at time t (Fig. 1a). The objects are represented as $\{O_1^t, O_2^t, O_3^t\}$, and a single force node F_1^t mediates the interaction. Here, the interaction scope controller $\rho_{\mathcal{O}}^t$ assigns high probability to the edges $J^t = \{E(O_1^t, F_1^t), E(O_2^t, F_1^t)\}$, indicating that both Ball 1 and Ball 2 contribute to generating the collision force. Similarly, the effect attribution controller $\rho_{\mathcal{O} \leftrightarrow \mathcal{F}}^t$ assigns high probability to the edges $I^t = \{E(F_1^t, O_1^{t+1}), E(F_1^t, O_2^{t+1})\}$, specifying that the force affects both balls post-collision. The resulting graph is illustrated in Fig. 1c. In contrast, when the balls are far apart and no interaction occurs, both controllers would output empty edge sets, resulting in a sparse graph with only self-connections during that time period (e.g., Ball 3 in Fig. 1c).

3.1.3. INDUCTIVE BIASES

To ground the flexible graph construction in realistic physical principles and mitigate the risk of overfitting to spurious connections, we incorporate two key inductive biases that reflect common patterns in object interactions. **Pairwise Interactions** restrict each force node to connect to exactly two different object nodes. This corresponds to a domain-motivated inductive bias that in many contact-like environments interactions are predominantly pairwise; it also keeps the controller’s action space tractable. This restriction can be lifted later to generalize to hypergraphs for more complicated systems (e.g., 3-body problems). We model a **mirroring (consistency) constraint**: if a force is generated by a set of objects, it can only be attributed to (i.e., affect) objects within that set. This encourages physically plausible and interpretable graphs by preventing attribution to unrelated objects. Importantly, this constraint does not enforce equal or symmetric effects; the effect attribution controller can attribute a force to one or both participants (Sec. 4.2), and the object update $f_{\mathcal{O}}$ can learn asymmetric influence magnitudes, including effectively zeroing influence on one side in domains with asymmetric interactions. Because $f_{\mathcal{F}}$ and $f_{\mathcal{O}}$ in Eq. 1 are invariant to the number of parents, CPM’s update mechanism itself does not require pairwise interactions; extending

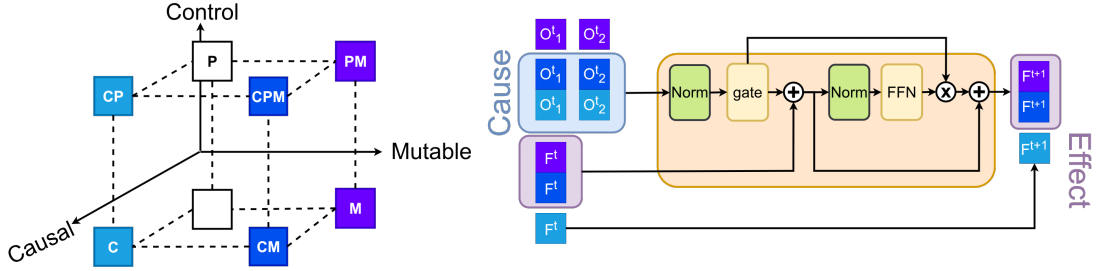


Figure 3: **Causal Process Block** (illustrating f_F): Modified transformer with attention replaced by gate mechanism (see Appendix C.3). Incoming causes O_i^t are pre-selected by the causal controllers. Latent vectors are divided into causal (C), control (P), and mutable (M) regions, enforcing structured updates.

the controllers to k -ary selections is a natural generalization, although we do not evaluate it here. In physical collisions, forces come in equal and opposite pairs acting on both objects. This design ensures physical consistency while maintaining computational efficiency. In environments with asymmetric interactions (e.g., large objects unaffected by small ones), the learned weights in f_O can effectively set the influence to zero (see Appendix C.1 for formal definitions).

4. Model

We base our model implementation on the Contrastively-trained Structured World Model (C-SWM; Kipf et al., 2020). The model consists of an *object-centric vision encoder*, an *action encoder*, and a *transition function* (Fig. 2). We keep the structure of the vision and action encoders intact, but modify the transition function.

The *vision encoder* is a CNN-based object extractor E_{ext} , operating directly on images and outputting I feature maps. Each feature map $m_i^t = [E_{\text{ext}}(s^t)]_i$ acts as an object mask where $[\dots]_i$ is the selection of the i^{th} feature map. An MLP-based object encoder E_{enc} with shared weights across objects maps the flattened feature map m_i^t to object latent representation: $O_i^t = E_{\text{enc}}(m_i^t)$. Additionally, an MLP-based *action encoder* maps action a^t to force latent representation: $F^t = A(a^t)$. Next, we introduce our new transition function (Sec. 4.1) before detailing how to construct the causal graph on the fly using reinforcement learning (Sec. 4.2).

4.1. Causal Process Block

Our main innovation is the *Causal Process Block* as a neural network implementation of a CPM (Fig. 3). Before introducing the technical details, we need to address a key challenge: not all components of force and object representations play the same role in causal interactions.

4.1.1. STRUCTURED REPRESENTATIONS

Let us revisit the example of collision between two balls (Fig. 1a). Their masses affect momentum transfer (causally relevant) but remain unchanged during the collision (immutable). In contrast, their velocities are both causally relevant and mutable. Meanwhile, visual properties like color may change due to lighting, but don’t affect the collision dynamics (mutable but not causally relevant). To capture these distinctions and enable our model to learn interpretable encodings that naturally separate these different physical properties, we factorize our representations along three key binary

subspaces of *Causal Relevance* (C), *Control Relevance* (P), and *Mutability* (M). The binary nature of each subspace arises from selective routing within the CPM (Fig. 3).

Causal Relevance (C) describes whether a component influences the dynamics of other objects. For instance mass and velocity affect collision outcomes ($C = 1$), but color does not ($C = 0$). Control Relevance (P) encodes whether a component is used by the control/policy functions for decision-making. For example, a controller deciding which balls are about to collide will rely on current positions and velocities ($P = 1$), but will ignore other properties such as mass or purely visual features like color ($P = 0$). Mutability (M) captures whether a component can change over time through interactions. For instance, an object being struck may change velocity ($M = 1$), while its mass remains constant ($M = 0$). Importantly, while the routing structure is fixed (based on $C/P/M$), the model learns end-to-end which features populate each subspace. We view this as an inductive bias that encourages interpretable factorization; it may be less suitable in domains where these axes are misaligned with the true generative factors, which we note as a limitation and a direction for future work.

4.1.2. TECHNICAL IMPLEMENTATION

We use two feedforward neural networks, $f_F(\dots; \theta_F)$ and $f_O(\dots; \theta_O)$, shared by all the force and object nodes respectively. The force vector $F_j^t := \bigoplus_{C,P,M \in \{1,0\}, (C,P,M) \neq (0,0,0)} F_j^{t,CPM}$ is the concatenation of all combinations of the C , P , M dimensions except for $(C, P, M) = (0, 0, 0)$, which must be omitted since forces, by definition, do not contain subspaces that are irrelevant to the causal process. This results in $2^3 - 1 = 7$ sub-vectors of equal size d_F , where a sub-vector’s identity determines how it is processed by the neural networks: the object-update function f_O and force-update function f_F operate exclusively on the causally relevant subspace ($C = 1$), while mutable parts are updated and immutable parts are copied unchanged ($M = 0$; Fig. 3). The object vector O_i^t is more straightforwardly divided into $2^3 = 8$ subvectors, i.e., all possible combinations of the C , P , M dimensions, including the $(C, P, M) = (0, 0, 0)$ subspace: $O_i^t := \bigoplus_{C,P,M \in \{1,0\}} O_i^{t,CPM}$. The extra subspace is present here for the network to learn to shift visual input features that are irrelevant to the causal process into this subspace, for example the object’s color in a collision event (See Appendices C.2 and C.3 for details). Note that the implementation of the causal process blocks $f_O(\dots | \theta_O)$ and $f_F(\dots | \theta_F)$ is similar to that of transformer blocks, but with the attention mechanism (Vaswani et al., 2017) replaced by indices of the chosen force and object nodes (tokens in transformers; see Appendix C.3 and Fig. 3 for more details). Unlike the attention mechanism of the transformer, I^t, J^t can also be an empty set. This is analogous to transformer attention assigning zero weight to all the tokens, which they cannot do by design.

4.2. Causal Controller

The main proposal of our model is that we perform graph construction through sequential decision making using the interaction scope and effect attribution controllers. We treat causal discovery as a multi-agent RL problem. One agent (the interaction scope policy $\pi_O(\mathcal{O}^t) := \rho_{\mathcal{O}}^t$) determines the scope of interacting objects, and another (the effect attribution policy $\pi_{O \leftrightarrow F}(\mathcal{O}^t, \mathcal{F}^{t+1}) := \rho_{\mathcal{O} \leftrightarrow \mathcal{F}}^t$) determines how force effects are attributed.

More specifically, the chosen indices I^t and J^t are provided by the agents $\pi_{O \leftrightarrow F}$ and π_O . The two agents alternate outputting an action. An action taken by π_O corresponds to two edge additions $E(O_i^t, F^{t+1}), E(O_j^t, F^{t+1}), i \neq j$ to the graph (selecting a pair of objects for interaction).

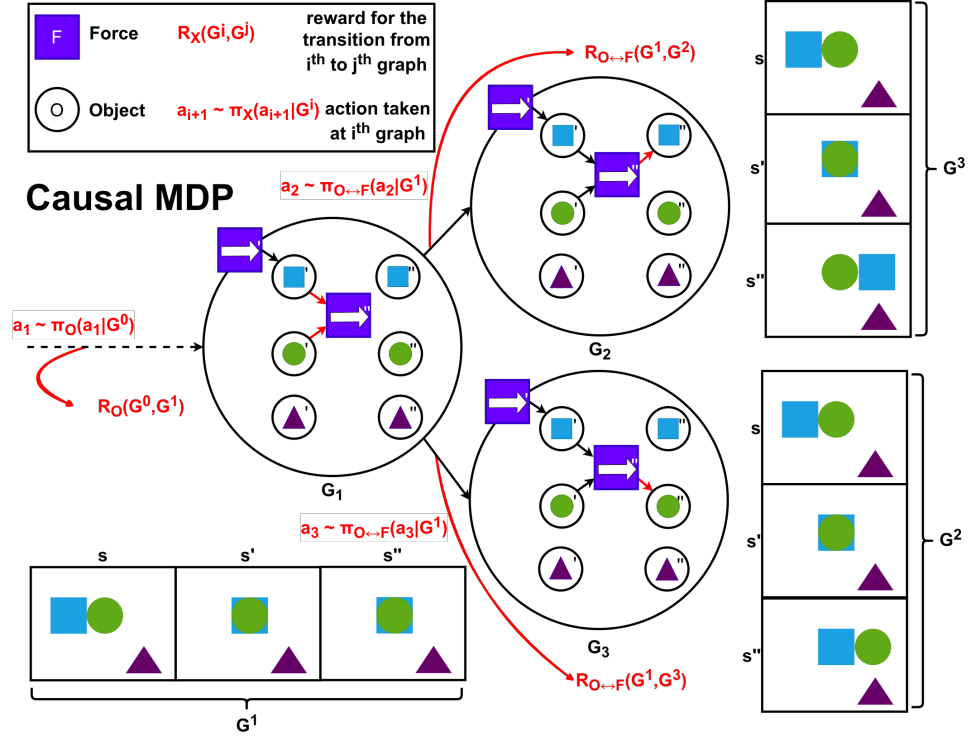


Figure 4: **Causal MDP** used by the reactive agents to construct causal process graphs. Agents successively add edges to the causal graph. Each causal graph hypothesis corresponds to a potential sequence of frames.

Whereas an action taken by $\pi_{O \leftrightarrow F}$ results in either one or two edge additions $E(F^t, O_i^t)$, $E(F^t, O_j^t)$ (attributing the force effect to either one or both objects; see Fig. 4). The index set I^t is sampled using the policy of the agent $\pi_{O \leftrightarrow F}$. Note that the policies only utilize the Control Relevant (P) features of the latent representations:

$$I^t \sim \pi_{O \leftrightarrow F}(I^t | G^t, W_O, W_F) = \text{softmax}(Q_{O \leftrightarrow F}(G^t, I^t | W_O, W_F)) \\ = \text{softmax}\left(\frac{(F^{t,C1M} W_F) \left([O_i^{t,C1M}; O_j^{t,C1M}; (O_i^{t,C1M} + O_j^{t,C1M})/2] W_O \right)^T}{d}\right), \quad (2)$$

where $W_O \in \mathbb{R}^{4d_O \times d}$, $W_F \in \mathbb{R}^{4d_F \times d}$, and $Q_{O \leftrightarrow F}$ is the corresponding Q-value. J^t , on the other hand, is sampled using the policy of the agent π_O :

$$J^t \sim \pi_O(J^t | G^t, W_{\tilde{O}}) = \sigma(Q_O(G^t, J^t | G^t, W_{\tilde{O}})) = \sigma((O_i^{t,C1M} W_{\tilde{O}}) \cdot (O_j^{t,C1M} W_{\tilde{O}})), \quad (3)$$

where $W_{\tilde{O}} \in \mathbb{R}^{4d_O \times d}$, σ is the sigmoid function, and Q_O is the corresponding Q-value.

We then define separate reward functions for $\pi_{O \leftrightarrow F}$ and π_O , modeled by MLPs parameterized by $\theta_{R_{O \leftrightarrow F}}$ and θ_{R_O} respectively:

$$R_{O \leftrightarrow F}(G^t, G^{t+1} | \theta_{R_{O \leftrightarrow F}}) = \text{MLP}\left(G_V^t, \mathbb{1}_{E(F^t, O_i^t)}, \mathbb{1}_{E(F^t, O_j^t)}, G_V^{t+1} | \theta_{R_{O \leftrightarrow F}}\right), \\ R_O(G^t, G^{t+1} | \theta_{R_O}) = \text{MLP}\left(G_V^t, \mathbb{1}_{E(O_i^t, F^{t+1}) \wedge E(O_j^t, F^{t+1})}, G_V^{t+1} | \theta_{R_O}\right), \quad (4)$$

where $G^t := (G_V^t, G_E^t)$ is the graph at time t and $\mathbb{1}_{E(\cdot, \cdot)}$ indicates the presence or absence of the edge $E(\cdot, \cdot)$. These reward functions are learned through inverse reinforcement learning (Ng and Russell, 2000), where the goal is to find a reward function whose corresponding optimal policy would select causal edges that would minimize prediction error.

4.3. Training Procedure

Edge selection and representation learning are tightly coupled: controllers require informative representations to choose edges, while learning informative representations depends on selecting reasonable edges. We therefore adopt a staged, EM-like alternating optimization (Dempster et al., 1977) to stabilize learning (Stage 1 warm-starts representation/dynamics under fixed controllers; Stages 2–3 alternate reward learning and policy improvement). We choose discrete, all-or-nothing edges to preserve sparse computation and avoid reverting to dense message passing during training. The overall goal is to learn a predictive world model (CPM) whose structure is determined by the policies $(\pi_O, \pi_{O \leftrightarrow F})$. This requires optimizing both the model parameters Θ and the policy parameters $\Psi := [W_O, W_F, W_{\tilde{O}}]$.

1. Prediction. At this stage, we freeze all the weights except $\Theta = [\theta_V; \theta_A; \theta_O; \theta_F]$ and sample edges $\{I^\tau, J^\tau\}_\tau$ using frozen controllers. This allows the vision encoder, action encoder, and transition functions (f_O, f_F) to learn useful representations before the controllers get the chance to optimize their behavior on these representations. We train the model parameters Θ using contrastive loss (Kipf et al., 2020):

$$\mathcal{L}_{\text{pred}}(\Theta | \beta, \mathcal{D}_{\text{pred}}) = \left\| \text{CPM} \left(V(s^t | \theta_V), A(a^t | \theta_A), (I^\tau, J^\tau)_{t^\tau} \middle| \theta_O, \theta_F \right) - V(s^{t+1} | \theta_V) \right\| + \max(0, \beta - \|V(\tilde{s}^t | \theta_V) - V(s^{t+1} | \theta_V)\|) \quad (5)$$

where $\mathcal{D}_{\text{pred}} = \{(s^t, a^t, s^{t+1}), \tilde{s}^t, \{(G^{t^\tau}, I^{t^\tau}, J^{t^\tau}, G^{t^\tau+1})\}_\tau\}_t$, V and A are the vision and action encoders, $\tilde{s}^t$ is a negative example sampled from the experience buffer, and the hinge margin β is set to 1 (following Kipf et al., 2020).

2. Expectation. In the second stage, we freeze all the weights but $\Theta_R = [\theta_{R_{O \leftrightarrow F}}; \theta_{R_O}]$ and use temporal difference (TD) loss (Watkins and Dayan, 1992) to learn reward functions:

$$\mathcal{L}_{\text{TD}}(\Theta_R | \mathcal{D}_{\text{TD}}, \Psi) = \sum_{X \in \{O \leftrightarrow F, O\}} \sum_\tau \left\| \underbrace{R_X(G^\tau, G^{\tau+1} | \Theta_{R_X})}_{\text{learned}} - \underbrace{\left(Q_X(G^\tau, I_X^\tau | \Psi) - \gamma \max_{I_X^{\tau+1}} Q_X(G^{\tau+1}, I_X^{\tau+1} | \Psi) \right)}_{\text{target}} \right\| \quad (6)$$

where $\mathcal{D}_{\text{TD}} = \{(G^\tau, I_O^\tau, I_{O \leftrightarrow F}^\tau, G^{\tau+1})\}_\tau$ and we set $\gamma = 0.9$.

3. Maximization. During the third stage, we freeze all but the policy parameters Ψ and use the same TD loss from above with target and learned terms reversed. The agents learn to select edges based on the rewards $R_{O \leftrightarrow F}$ and R_O .

The overall optimization landscape is complex (see Appendix D for further optimization details); the optimum represents a state where the agents select the sparse causal graph that yields the minimal prediction loss for the CPM, and simultaneously, the reward MLPs stabilize under the IRL objective.

5. Experiments

We hypothesize that our model outperforms models that assume dense causal graphs to capture physical interactions in: 1) longer prediction horizons; 2) test-time generalization across unobservable

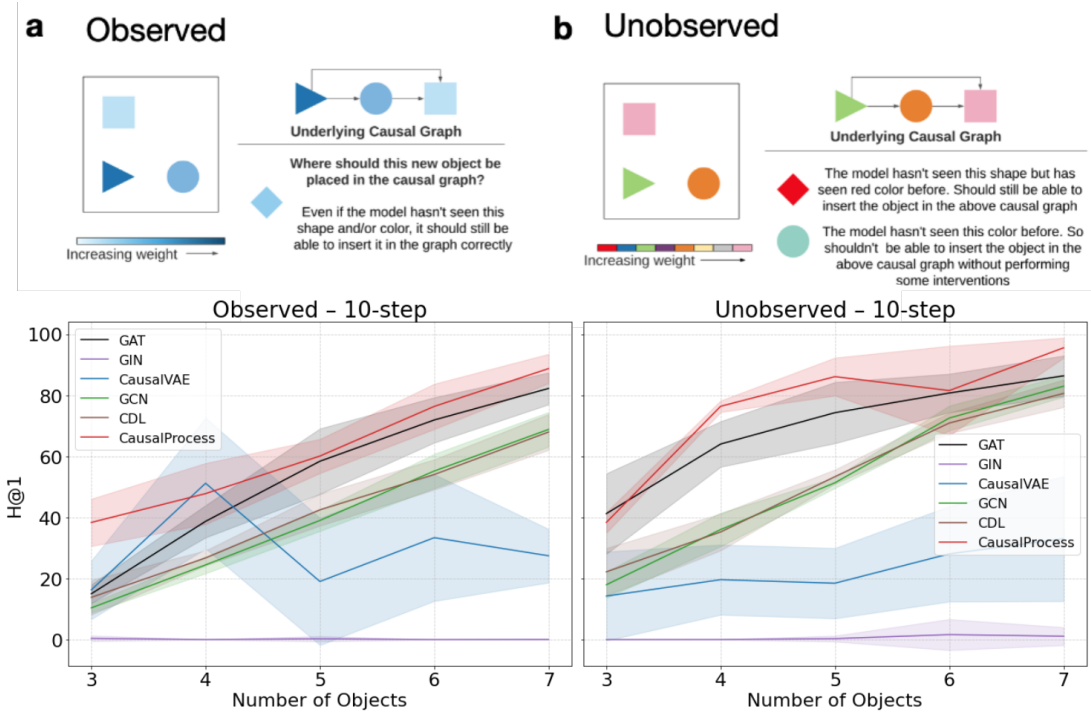


Figure 5: **Prediction results for a synthetic physics environment** in a) observed and b) unobserved settings (Ke et al., 2021). **Top:** Description of the task. **Bottom:** Prediction metric vs number of objects after 10 steps (average of 10 seeds). For remaining baselines, see Fig. 7 in Appendix E.

properties; 3) robustness with regards to the number of objects in the scene; 4) solving downstream tasks. We use the *physics environment* designed by Ke et al. (2021) to empirically answer these questions (Fig. 5 top). The environment consists of different objects colored according to their weights. The only force in this environment is pushing (double-pushes are not allowed) and only heavier objects can push lighter ones. The environment has two settings: an *observed* setting (Fig. 5a) where weight corresponds to the intensity of a particular color and an *unobserved* setting (Fig. 5b) where different colors do not systematically map to different weights.

Benchmark choice and empirical scope. We adopt the synthetic physics benchmark of Ke et al. (2021) because it is an established testbed for visual causal discovery in model-based RL and enables direct comparison to a broad set of prior baselines under a standardized protocol. The benchmark is intentionally controlled: interactions are sparse and pairwise, observations are strongly object-centric, and the only explicit interaction is a pushing force (with additional constraints such as no double-pushes and asymmetric pushing between heavy/light objects). These restrictions allow us to isolate the main claim of CPMs, that dynamically constructing *sparse, time-varying* causal graphs improves long-horizon prediction and downstream planning, while keeping confounds minimal. We view evaluation in more diverse physical domains as an important next step (see Sec. 6).

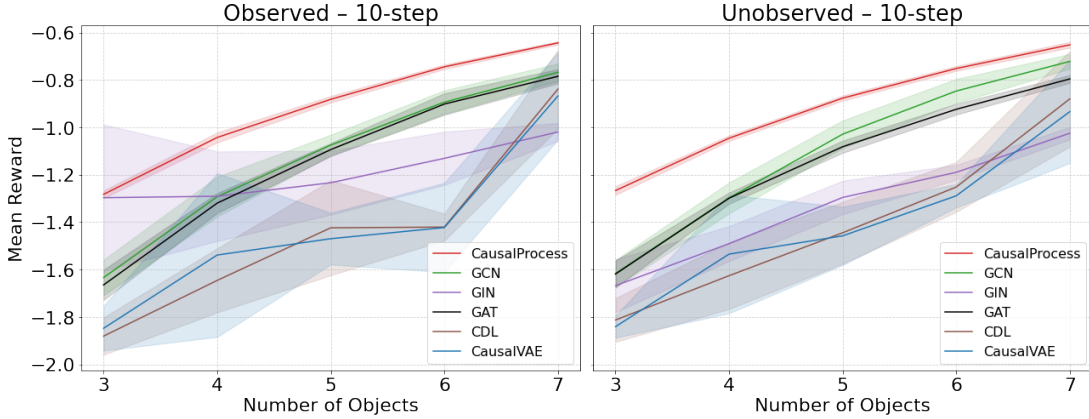


Figure 6: **Downstream RL results** over number of objects. Mean reward vs number of objects. All results are the average of 10 seeds. For remaining baselines, see Fig. 8 in Appendix E.

5.1. Comparison Baselines

We compare our model against 5 causal and graph neural network baselines: a graph attention network (GAT) (Veličković et al., 2018), a graph isomorphism network (GIN) (Xu et al., 2019), a causal variational auto-encoder (CausalVAE) (Yang et al., 2021), a graph convolutional network (GCN) (Kipf and Welling, 2017), and a causal dynamics learning network (CDL) (Wang et al., 2022). Additionally, we compare the model against the 5 baselines from Ke et al. (2021): a graph neural network (GNN) (Scarselli et al., 2009), a transformer network (Vaswani et al., 2017), a recurrent independent mechanisms (RIM) network (Goyal et al., 2021b), a schema / object-file factorization network (SCOFF) (Goyal et al., 2021a), and a modular network that has a separate MLP to model each object’s dynamics. We do not include RL-based tabular causal discovery methods as experimental baselines because they output a single global graph over pre-defined variables and do not define a visual world-model predictor under our evaluation metrics (Zhu et al., 2020; Sauter et al., 2024; Duong et al., 2025).

5.2. Prediction Metrics

To investigate robustness towards the length of prediction horizons, we trained the model to make 1-step predictions in the *Observed* setting with several objects and then tested for 5 (Fig. 9) and 10 steps (Fig. 5a bottom, Fig. 7). We used Hits at Rank 1 (H@1) to measure model performance as an all-or-nothing metric measuring how often the rank of the predicted representation was 1 when ranked against all reference state representations. Here, our model broadly outperformed the baseline models, with the gap increasing over longer time horizons.

Next, to estimate the test-time generalization across unobservable properties, we trained our model in the *Unobserved* setting where generalization at test time is harder due to previously unseen colors (weights). Again, our model broadly outperformed the baselines displaying capacity to generalize also in this domain (Fig. 5b bottom; see Fig. 7 and Fig. 9 for more results).

5.3. Downstream RL tasks

To make sure the above metrics overlap with the learned model’s usefulness for downstream tasks, we also tested our CPM’s capacity to serve as a world model for a model-based RL agent. The agent’s task was to move an object to a certain location each taken step resulting in negative reward.

In both Observed and Unobserved settings, the agent with CPM as model of the environment broadly outperformed the baselines for all objects in 10-step unrolling of the learned model (Fig. 6, Fig. 8).

5.4. Ablations and Graph Recovery

We include a component ablation study in Appendix F. We additionally evaluate (i) graph recovery against the simulator’s ground-truth interaction graph and (ii) the semantics of the learned (C, P, M) subspaces via linear probing; see Appendix G.

6. Discussion

In this paper, we introduced the Causal Process Framework (CPF) as a novel approach for modeling the dynamics of physical object interactions. Our key contribution is the Causal Process Model (CPM), which implements this framework by treating the edge distributions inherent to CPF as a reinforcement learning policy. Instead of the soft, dense connections typical of many baselines (Veličković et al., 2018; Goyal et al., 2021a; Xu et al., 2019; Yang et al., 2021; Kipf and Welling, 2017; Wang et al., 2022; Scarselli et al., 2009; Vaswani et al., 2017; Goyal et al., 2021b), our model employs RL agents to dynamically construct sparse, time-varying causal graphs. Our experiments in a simulated physics environment (Ke et al., 2021) show that this approach not only improves prediction accuracy and downstream task performance compared to baselines, but also excels in generalization and scalability.

The superior performance of our model, particularly over longer prediction horizons and with a varying number of objects, lends strong support to our central hypothesis. We argue that by explicitly modeling only active causal links, the CPM avoids the pitfalls of dense message-passing architectures (Barbero et al., 2024a; Alon and Yahav, 2021; Barbero et al., 2024b; Giovanni et al., 2023, 2024; Topping et al., 2022; Scarselli et al., 2009; Battaglia et al., 2018). Our discrete, “all-or-nothing” connections, determined by a goal-oriented RL agent, preserve the salience of individual interactions. This leads to more robust and precise world models, which proved crucial for the model-based RL agent’s success in downstream tasks. Furthermore, the model’s ability to generalize to unobserved object properties suggests that it learns an underlying model of physical dynamics rather than memorizing superficial correlations.

Despite these promising results, the present work has several limitations that open clear avenues for future research. A primary limitation is that our empirical evaluation is conducted in a single benchmark environment with strongly object-centric visuals, pairwise interactions, and a single force type (pushing). While this setting provides a clean comparison against established baselines, it does not by itself demonstrate robustness to substantially more complex physical interactions or perception. Importantly, these restrictions are not fundamental to the Causal Process Framework / CPM. Our current pairwise interaction constraint is a deliberate inductive bias (Sec. 3.1.3) and can be lifted to support hypergraph-style interactions (e.g., three-body effects) by allowing force nodes to connect to sets of objects rather than pairs. The object-centric structure in our experiments reflects the chosen vision front-end (Sec. 4) rather than a requirement of CPMs. An important direction is to pair CPMs with more general perception modules that handle clutter, occlusion, and imperfect segmentation. Additionally, we do not empirically compare against continuous edge-selection relaxations (e.g., Concrete/Gumbel masking) (Jang et al., 2017; Louizos et al., 2018; Maddison et al., 2017) and we do not provide an extensive hyperparameter sensitivity analysis; we leave these evaluations to future work.

Acknowledgments

Funded by the Deutsche Forschungsgemeinschaft (DFG, German Research Foundation) under Germany’s Excellence Strategy – EXC number 2064/1 – Project number 390727645. C.M.W is supported by the European Research Council (ERC) under the European Union’s Horizon 2020 research and innovation programme (C4: 101164709), the Deutsche Forschungsgemeinschaft (German Research Foundation, DFG) under Germany’s Excellence Strategy (EXC 3066/1 “The Adaptive Mind”, Project No. 533717223), and the Excellence Cluster “Reasonable AI” by the Deutsche Forschungsgemeinschaft (German Research Foundation, DFG) under Germany’s Excellence Strategy – EXC-3057. The authors thank the International Max Planck Research School for Intelligent Systems (IMPRS-IS) for supporting Turan Orujlu. The authors also thank Antonio Orvieto, Siyuan Guo, and Patrik Reizinger for insightful discussions.

References

- Uri Alon and Eran Yahav. On the bottleneck of graph neural networks and its practical implications. In *9th International Conference on Learning Representations, ICLR 2021, Virtual Event, Austria, May 3-7, 2021*. OpenReview.net, 2021. URL <https://openreview.net/forum?id=i80OPhOCVH2>.
- Federico Barbero, Andrea Banino, Steven Kapturowski, Dharshan Kumaran, João G. M. Araújo, Alex Vitvitskiy, Razvan Pascanu, and Petar Velickovic. Transformers need glasses! information over-squashing in language tasks. *CoRR*, abs/2406.04267, 2024a. doi: 10.48550/ARXIV.2406.04267. URL <https://doi.org/10.48550/arXiv.2406.04267>.
- Federico Barbero, Ameya Velingker, Amin Saberi, Michael M. Bronstein, and Francesco Di Giovanni. Locality-aware graph rewiring in gnns. In *The Twelfth International Conference on Learning Representations, ICLR 2024, Vienna, Austria, May 7-11, 2024*. OpenReview.net, 2024b. URL <https://openreview.net/forum?id=4Ua4hKiAJX>.
- Elias Bareinboim, S Lee, and J Zhang. An introduction to causal reinforcement learning, 2021.
- Peter W. Battaglia, Jessica B. Hamrick, Victor Bapst, Alvaro Sanchez-Gonzalez, Vinícius Flores Zambaldi, Mateusz Malinowski, Andrea Tacchetti, David Raposo, Adam Santoro, Ryan Faulkner, Çağlar Gülçehre, H. Francis Song, Andrew J. Ballard, Justin Gilmer, George E. Dahl, Ashish Vaswani, Kelsey R. Allen, Charles Nash, Victoria Langston, Chris Dyer, Nicolas Heess, Daan Wierstra, Pushmeet Kohli, Matthew M. Botvinick, Oriol Vinyals, Yujia Li, and Razvan Pascanu. Relational inductive biases, deep learning, and graph networks. *CoRR*, abs/1806.01261, 2018. URL <http://arxiv.org/abs/1806.01261>.
- Tineke Blom, Stephan Bongers, and Joris M. Mooij. Beyond structural causal models: Causal constraints models. In Ryan P. Adams and Vibhav Gogate, editors, *Proceedings of The 35th Uncertainty in Artificial Intelligence Conference*, volume 115 of *Proceedings of Machine Learning Research*, pages 585–594. PMLR, 22–25 Jul 2020. URL <https://proceedings.mlr.press/v115/blom20a.html>.
- Philip A. Boeken and Joris M. Mooij. Dynamic structural causal models. 2024. URL <https://api.semanticscholar.org/CorpusID:270213849>.

- Lars Buesing, Theophane Weber, Yori Zwols, Nicolas Heess, Sébastien Racanière, Arthur Guez, and Jean-Baptiste Lespiau. Woulda, coulda, shoulda: Counterfactually-guided policy search. In *7th International Conference on Learning Representations, ICLR 2019, New Orleans, LA, USA, May 6-9, 2019*. OpenReview.net, 2019. URL <https://openreview.net/forum?id=BJG0voC9YQ>.
- Martin V. Butz. Which structures are out there? learning predictive compositional concepts based on social sensorimotor explorations, 2017. ISSN 978-3-95857-138-9.
- 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. ISSN 0149-7634. doi: <https://doi.org/10.1016/j.neubiorev.2024.105948>. URL <https://www.sciencedirect.com/science/article/pii/S0149763424004172>.
- Arthur P. Dempster, Nan M. Laird, and Donald B. Rubin. Maximum likelihood from incomplete data via the EM algorithm. *Journal of the Royal Statistical Society. Series B (Methodological)*, 39(1):1–38, 1977. doi: 10.1111/j.2517-6161.1977.tb01600.x. URL <http://www.jstor.org/stable/2984875>. JSTOR: 2984875.
- Hanna M Dettki, Brenden M Lake, Charley M Wu, and Bob Rehder. Do large language models reason causally like us? even better? *arXiv preprint arXiv:2502.10215*, 2025.
- Phil Dowe. *Physical Causation*. New York, 2000.
- Bao Duong, Hung Le, Biwei Huang, and Thin Nguyen. Reinforcement learning for causal discovery without acyclicity constraints. *Trans. Mach. Learn. Res.*, 2025, 2025. URL <https://openreview.net/forum?id=sNzBi8rZTy>.
- Maxime Gasse, Damien Grasset, Guillaume Gaudron, and Pierre-Yves Oudeyer. Using confounded data in latent model-based reinforcement learning. *Trans. Mach. Learn. Res.*, 2023, 2023. URL <https://openreview.net/forum?id=nFWRuJXPkU>.
- Tobias Gerstenberg, Noah D. Goodman, David A. Lagnado, and Joshua B. Tenenbaum. A counterfactual simulation model of causal judgments for physical events. *Psychological review*, 2020. URL <https://api.semanticscholar.org/CorpusID:235361720>.
- Francesco Di Giovanni, Lorenzo Giusti, Federico Barbero, Giulia Luise, Pietro Lio, and Michael M. Bronstein. On over-squashing in message passing neural networks: The impact of width, depth, and topology. In Andreas Krause, Emma Brunskill, Kyunghyun Cho, Barbara Engelhardt, Sivan Sabato, and Jonathan Scarlett, editors, *International Conference on Machine Learning, ICML 2023, 23-29 July 2023, Honolulu, Hawaii, USA*, volume 202 of *Proceedings of Machine Learning Research*, pages 7865–7885. PMLR, 2023. URL <https://proceedings.mlr.press/v202/di-giovanni23a.html>.
- Francesco Di Giovanni, T. Konstantin Rusch, Michael M. Bronstein, Andreea Deac, Marc Lackenby, Siddhartha Mishra, and Petar Velickovic. How does over-squashing affect the power of gnns? *Trans. Mach. Learn. Res.*, 2024, 2024. URL <https://openreview.net/forum?id=KJRoQvRWNs>.

- Anirudh Goyal, Alex Lamb, Phanideep Gampa, Philippe Beaudoin, Charles Blundell, Sergey Levine, Yoshua Bengio, and Michael Curtis Mozer. Factorizing declarative and procedural knowledge in structured, dynamical environments. In *International Conference on Learning Representations*, 2021a. URL <https://openreview.net/forum?id=VVdmjgu7pKM>.
- Anirudh Goyal, Alex Lamb, Jordan Hoffmann, Shagun Sodhani, Sergey Levine, Yoshua Bengio, and Bernhard Schölkopf. Recurrent independent mechanisms. In *9th International Conference on Learning Representations, ICLR 2021, Virtual Event, Austria, May 3-7, 2021*. OpenReview.net, 2021b. URL <https://openreview.net/forum?id=mLcmdlEUxy->.
- Christian Gumbsch, Martin V Butz, and Georg Martius. Sparsely changing latent states for prediction and planning in partially observable domains. *Advances in Neural Information Processing Systems*, 34:17518–17531, 2021.
- Eric Jang, Shixiang Gu, and Ben Poole. Categorical reparameterization with gumbel-softmax. In *5th International Conference on Learning Representations, ICLR 2017, Toulon, France, April 24-26, 2017, Conference Track Proceedings*. OpenReview.net, 2017. URL <https://openreview.net/forum?id=rkE3y85ee>.
- Nan Rosemary Ke, Aniket Didolkar, Sarthak Mittal, Anirudh Goyal, Guillaume Lajoie, Stefan Bauer, Danilo Jimenez Rezende, Michael Mozer, Yoshua Bengio, and Chris Pal. Systematic evaluation of causal discovery in visual model based reinforcement learning. In Joaquin Vanschoren and Sai-Kit Yeung, editors, *Proceedings of the Neural Information Processing Systems Track on Datasets and Benchmarks 1, NeurIPS Datasets and Benchmarks 2021, December 2021, virtual*, 2021. URL <https://datasets-benchmarks-proceedings.neurips.cc/paper/2021/hash/8f121ce07d74717e0b1f21d122e04521-Abstract-round2.html>.
- Thomas N. Kipf and Max Welling. Semi-supervised classification with graph convolutional networks. In *International Conference on Learning Representations*, 2017. URL <https://openreview.net/forum?id=SJU4ayYgl>.
- Thomas N. Kipf, Elise van der Pol, and Max Welling. Contrastive learning of structured world models. In *8th International Conference on Learning Representations, ICLR 2020, Addis Ababa, Ethiopia, April 26-30, 2020*. OpenReview.net, 2020. URL <https://openreview.net/forum?id=Hlgax6VtDB>.
- Richard D. Lange and Konrad P. Kording. Causality in the human niche: lessons for machine learning. *CoRR*, abs/2506.13803, 2025. doi: 10.48550/ARXIV.2506.13803. URL <https://doi.org/10.48550/arXiv.2506.13803>.
- Anson Lei, Bernhard Schölkopf, and Ingmar Posner. Spartan: A sparse transformer learning local causation. *arXiv preprint arXiv:2411.06890*, 2024.
- Christos Louizos, Max Welling, and Diederik P. Kingma. Learning sparse neural networks through l₀ regularization. In *6th International Conference on Learning Representations, ICLR 2018, Vancouver, BC, Canada, April 30 - May 3, 2018, Conference Track Proceedings*. OpenReview.net, 2018. URL <https://openreview.net/forum?id=H1Y8hhg0b>.

- Chris J. Maddison, Andriy Mnih, and Yee Whye Teh. The concrete distribution: A continuous relaxation of discrete random variables. In *5th International Conference on Learning Representations, ICLR 2017, Toulon, France, April 24-26, 2017, Conference Track Proceedings*. OpenReview.net, 2017. URL <https://openreview.net/forum?id=S1jE5L5gl>.
- Valentyn Melnychuk, Dennis Frauen, and Stefan Feuerriegel. Causal transformer for estimating counterfactual outcomes. In Kamalika Chaudhuri, Stefanie Jegelka, Le Song, Csaba Szepesvári, Gang Niu, and Sivan Sabato, editors, *International Conference on Machine Learning, ICML 2022, 17-23 July 2022, Baltimore, Maryland, USA*, volume 162 of *Proceedings of Machine Learning Research*, pages 15293–15329. PMLR, 2022. URL <https://proceedings.mlr.press/v162/melnichuk22a.html>.
- Andrew Y. Ng and Stuart Russell. Algorithms for inverse reinforcement learning. In Pat Langley, editor, *Proceedings of the Seventeenth International Conference on Machine Learning (ICML 2000), Stanford University, Stanford, CA, USA, June 29 - July 2, 2000*, pages 663–670. Morgan Kaufmann, 2000.
- Eshaan Nichani, Alex Damian, and Jason D. Lee. How transformers learn causal structure with gradient descent. In *Forty-first International Conference on Machine Learning, ICML 2024, Vienna, Austria, July 21-27, 2024*. OpenReview.net, 2024. URL <https://openreview.net/forum?id=jNM4imlHZv>.
- Judea Pearl. *Causality: Models, Reasoning and Inference*. Cambridge University Press, USA, 2nd edition, 2009. ISBN 052189560X.
- Silviu Pitis, Elliot Creager, and Animesh Garg. Counterfactual data augmentation using locally factored dynamics. *Advances in Neural Information Processing Systems*, 33:3976–3990, 2020.
- B. T. Polyak and A. B. Juditsky. Acceleration of stochastic approximation by averaging. *SIAM Journal on Control and Optimization*, 30(4):838–855, 1992. doi: 10.1137/0330046. URL <https://doi.org/10.1137/0330046>.
- James Robins and Miguel Hernan. *Estimation of the causal effects of time-varying exposure*, pages 553–599. 08 2008. ISBN 978-1-58488-658-7. doi: 10.1201/9781420011579.ch23.
- Raanan Y. Rohekar, Yaniv Gurwicz, and Shami Nisimov. Causal interpretation of self-attention in pre-trained transformers. In Alice Oh, Tristan Naumann, Amir Globerson, Kate Saenko, Moritz Hardt, and Sergey Levine, editors, *Advances in Neural Information Processing Systems 36: Annual Conference on Neural Information Processing Systems 2023, NeurIPS 2023, New Orleans, LA, USA, December 10 - 16, 2023*, 2023. URL http://papers.nips.cc/paper_files/paper/2023/hash/642a321fba8a0f03765318e629cb93ea-Abstract-Conference.html.
- Paul K. Rubenstein, Stephan Bongers, Joris M. Mooij, and Bernhard Scholkopf. From deterministic odes to dynamic structural causal models. In *Conference on Uncertainty in Artificial Intelligence*, 2016. URL <https://api.semanticscholar.org/CorpusID:2303037>.
- Donald B. Rubin. Bayesian Inference for Causal Effects: The Role of Randomization. *The Annals of Statistics*, 6(1):34 – 58, 1978. doi: 10.1214/aos/1176344064. URL <https://doi.org/10.1214/aos/1176344064>.

- Bertrand Russell. *Human Knowledge: Its Scope and Limits*. Routledge, London and New York, 1948.
- Wesley C. Salmon. *Scientific Explanation and the Causal Structure of the World*. Princeton University Press, 1984.
- Andreas Sauter, Nicolò Botteghi, Erman Acar, and Aske Plaat. CORE: towards scalable and efficient causal discovery with reinforcement learning. In Mehdi Dastani, Jaime Simão Sichman, Natasha Alechina, and Virginia Dignum, editors, *Proceedings of the 23rd International Conference on Autonomous Agents and Multiagent Systems, AAMAS 2024, Auckland, New Zealand, May 6-10, 2024*, pages 1664–1672. International Foundation for Autonomous Agents and Multiagent Systems / ACM, 2024. doi: 10.5555/3635637.3663027. URL <https://dl.acm.org/doi/10.5555/3635637.3663027>.
- Franco Scarselli, Marco Gori, Ah Chung Tsoi, Markus Hagenbuchner, and Gabriele Monfardini. The graph neural network model. *IEEE Trans. Neural Networks*, 20(1):61–80, 2009. doi: 10.1109/TNN.2008.2005605. URL <https://doi.org/10.1109/TNN.2008.2005605>.
- Bernhard Schölkopf, Francesco Locatello, Stefan Bauer, Nan Rosemary Ke, Nal Kalchbrenner, Anirudh Goyal, and Yoshua Bengio. Toward causal representation learning. *Proceedings of the IEEE*, 109(5):612–634, 2021.
- Maximilian Seitzer, Bernhard Schölkopf, and Georg Martius. Causal influence detection for improving efficiency in reinforcement learning. *Advances in Neural Information Processing Systems*, 34: 22905–22918, 2021.
- Xiao Shou, Debarun Bhattacharjya, Tian Gao, Dharmashankar Subramanian, Oktie Hassanzadeh, and Kristin P. Bennett. Pairwise causality guided transformers for event sequences. In Alice Oh, Tristan Naumann, Amir Globerson, Kate Saenko, Moritz Hardt, and Sergey Levine, editors, *Advances in Neural Information Processing Systems 36: Annual Conference on Neural Information Processing Systems 2023, NeurIPS 2023, New Orleans, LA, USA, December 10 - 16, 2023*, 2023. URL http://papers.nips.cc/paper_files/paper/2023/hash/91b047c5f5bd41ef56bfaf4ad0bd19e3-Abstract-Conference.html.
- Brian Skyrms. Causal necessity. *Philosophy of Science*, 48(2):329–335, 1981. doi: 10.1086/289003.
- Jake Topping, Francesco Di Giovanni, Benjamin Paul Chamberlain, Xiaowen Dong, and Michael M. Bronstein. Understanding over-squashing and bottlenecks on graphs via curvature. In *The Tenth International Conference on Learning Representations, ICLR 2022, Virtual Event, April 25-29, 2022*. OpenReview.net, 2022. URL <https://openreview.net/forum?id=7UmjRGzp-A>.
- Núria Armengol Urpí, Marco Bagatella, Marin Vlastelica, and Georg Martius. Causal action influence aware counterfactual data augmentation. *arXiv preprint arXiv:2405.18917*, 2024.
- Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N. Gomez, Lukasz Kaiser, and Illia Polosukhin. Attention is all you need. In Isabelle Guyon, Ulrike von Luxburg, Samy Bengio, Hanna M. Wallach, Rob Fergus, S. V. N. Vishwanathan, and Roman Garnett, editors, *Advances in Neural Information Processing Systems 30: Annual Conference on Neural Information Processing Systems 2017, December 4-9, 2017, Long Beach, CA, USA*, pages

- 5998–6008, 2017. URL <https://proceedings.neurips.cc/paper/2017/hash/3f5ee243547dee91fbd053c1c4a845aa-Abstract.html>.
- Petar Veličković, Guillem Cucurull, Arantxa Casanova, Adriana Romero, Pietro Liò, and Yoshua Bengio. Graph attention networks. In *International Conference on Learning Representations*, 2018. URL <https://openreview.net/forum?id=rJXMpikCZ>.
- Zizhao Wang, Xuesu Xiao, Zifan Xu, Yuke Zhu, and Peter Stone. Causal dynamics learning for task-independent state abstraction. *CoRR*, abs/2206.13452, 2022. URL <https://doi.org/10.48550/arXiv.2206.13452>.
- Christopher J. C. H. Watkins and Peter Dayan. Technical note q-learning. *Mach. Learn.*, 8:279–292, 1992. doi: 10.1007/BF00992698. URL <https://doi.org/10.1007/BF00992698>.
- Marcel Weber. On the incompatibility of dynamical biological mechanisms and causal graphs. *Philosophy of Science*, 83(5):959–971, 2016. ISSN 00318248, 1539767X. URL <https://www.jstor.org/stable/26551794>.
- Moritz Willig, Tim Nelson Tobiasch, Florian Peter Busch, Jonas Seng, Devendra Singh Dhami, and Kristian Kersting. Systems with switching causal relations: A meta-causal perspective. In *The Thirteenth International Conference on Learning Representations, ICLR 2025, Singapore, April 24-28, 2025*. OpenReview.net, 2025. URL <https://openreview.net/forum?id=J9VogDTa1W>.
- Kevin Xia, Kai-Zhan Lee, Yoshua Bengio, and Elias Bareinboim. The causal-neural connection: Expressiveness, learnability, and inference. In Marc’Aurelio Ranzato, Alina Beygelzimer, Yann N. Dauphin, Percy Liang, and Jennifer Wortman Vaughan, editors, *Advances in Neural Information Processing Systems 34: Annual Conference on Neural Information Processing Systems 2021, NeurIPS 2021, December 6-14, 2021, virtual*, pages 10823–10836, 2021. URL <https://proceedings.neurips.cc/paper/2021/hash/5989add1703e4b0480f75e2390739f34-Abstract.html>.
- Keyulu Xu, Weihua Hu, Jure Leskovec, and Stefanie Jegelka. How powerful are graph neural networks? In *International Conference on Learning Representations*, 2019. URL <https://openreview.net/forum?id=ryGs6iA5Km>.
- Mengyue Yang, Furui Liu, Zhitang Chen, Xinwei Shen, Jianye Hao, and Jun Wang. Causalvae: Disentangled representation learning via neural structural causal models. In *2021 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 9588–9597, 2021. doi: 10.1109/CVPR46437.2021.00947.
- Matej Zecevic, Devendra Singh Dhami, Petar Velickovic, and Kristian Kersting. Relating graph neural networks to structural causal models. *CoRR*, abs/2109.04173, 2021. URL <https://arxiv.org/abs/2109.04173>.
- Shengyu Zhu, Ignavier Ng, and Zhitang Chen. Causal discovery with reinforcement learning. In *8th International Conference on Learning Representations, ICLR 2020, Addis Ababa, Ethiopia, April 26-30, 2020*. OpenReview.net, 2020. URL <https://openreview.net/forum?id=S1g2skStPB>.

Appendix A. Object-centric Causal Dynamics

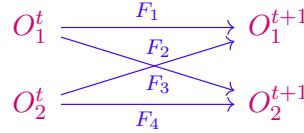
Consider two objects O_1 and O_2 (depicted in magenta) with a force F (depicted in violet) acting on them:

$$O_1 \xrightarrow{F} O_2$$

This recovers the familiar structure of a directed acyclic graphs (DAGs) from Pearl’s causal formalism Pearl (2009). However, in physical interactions, such as in a collision, it is not always clear which object is the “cause” since both are affected simultaneously. A more intuitive representation would be a bidirectional edge:

$$O_1 \xleftrightarrow{F} O_2$$

However, DAGs prohibit cycles and bidirectional edges. To resolve this, we introduce *temporal dynamics* which represent causal effects as unfolding over time rather than as a simultaneous influence. Thus, a collision between object O_1 and object O_2 yields forces F_2 and F_3 as emerging from the past state and influencing future object states:

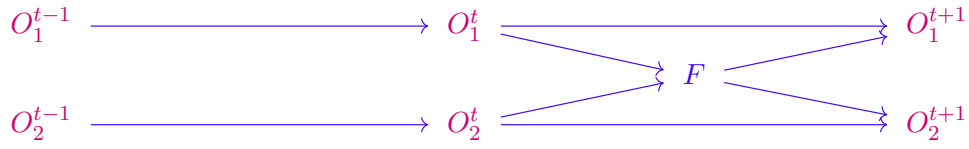


Yet, this representation still has drawbacks. Specifically, we break the identity of the force F into F_2 and F_3 which, in principle, can act as separate causal links (F_1 and F_4 can be thought of as inertia). This becomes apparent when interventions are applied. Let us imagine that somebody picks up object O_1 just before it collides with O_2 . This can be represented by a do-calculus-like intervention applied to either O_1^t or O_1^{t+1} :

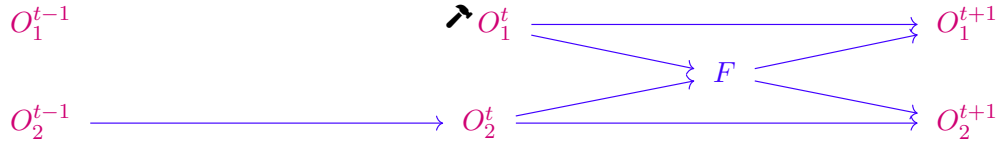


When the intervention is applied to O_1^t , the graph structure is preserved, thus implying a no-collision scenario (one of the balls was lifted). Yet, the same graph can also imply a collision scenario. This kind of setup necessitates having causal links between objects that can potentially collide irrespective of the actualization of said collision. While this approach can work in principle, it results in extremely dense graphs with complete subgraphs per time step, especially in cluttered scenes. Ideally, we would like to have causal links in our graph if there is an actualized interaction between the involved objects.

On the other hand, intervening on O_1^{t+1} results in a graph with counter-intuitive interpretation: O_1 gets lifted at time step $t + 1$, while O_2 behaves as if a collision has happened. This is due to the split of F into F_2 and F_3 since an intervention removes F_3 while leaving F_2 untouched. To tackle the aforementioned issues, let us re-imagine force edges as nodes and re-introduce F_2 and F_3 as a single node F and extend the time horizon by a step.



Now, imagine, just like before, the ball O_1 gets picked up at time step t . In do-calculus terms, this amounts to intervention to O_1^t which results in mutilation of the edge $O_1^{t-1} \rightarrow O_1^t$



While the problem of splitting of the force identity seems to be resolved here, the graph structure modeling the collision remains preserved despite the intervention. As mentioned before, this can be addressed by complete subgraphs per time step, which is not desirable for our purposes. This problem arises due to the inclusion of time dynamics into our graphs. Unlike in Pearlian Causality, in physics, interventions at a time step have implications for the causal connections corresponding to downstream time steps. To account for that, we have to re-imagine interventions under a new framework that takes physical processes and time into account (see Algorithm 1).

Appendix B. Algorithms

Algorithm 1: Interventions under Causal Process Framework

Input: $\mathcal{F}^t, \mathcal{O}^t, f_F, f_O, \rho_{\mathcal{O} \leftrightarrow \mathcal{F}}^t, \rho_{\mathcal{O}}^t, \text{do} \left(F_{*}^{\tilde{t}}, * \rightarrow \imath \right), G_{\mathcal{O}^1: \mathcal{O}^T}, t \in \{1, \dots, T\}, i \in \{1, \dots, I\}, j \in \{1, \dots, J\}$
Output: $\tilde{G}_{\mathcal{O}^1: \mathcal{O}^T}$
 $\tilde{\mathcal{O}}^{\tilde{t}-1} := \mathcal{O}^{\tilde{t}-1};$
 $\tilde{\mathcal{F}}^{\tilde{t}} := \{F_{*}^{\tilde{t}}\} \dot{\cup} \mathcal{F}^{\tilde{t}};$
 $J^{\tilde{t}} := \{(i, j)\} \text{ s.t. } E \left(\mathcal{O}_i^{\tilde{t}-1}, F_j^{\tilde{t}} \right) \in G_{\mathcal{O}^{\tilde{t}-1}, \mathcal{F}^{\tilde{t}}}^{\mathcal{E}};$
 $\tilde{G}_{\mathcal{O}^1: \mathcal{O}^{\tilde{t}-1}} := G_{\mathcal{O}^1: \mathcal{O}^{\tilde{t}-1}};$
for $t = \tilde{t}$ **to** T **do**
 $\tilde{G}_{\mathcal{O}^1: \mathcal{F}^t} := \left(\tilde{\mathcal{F}}^t \dot{\cup} \tilde{G}_{\mathcal{O}^1: \mathcal{O}^{t-1}}^{\mathcal{V}}, \left\{ E \left(\tilde{\mathcal{O}}_i^{t-1}, \tilde{F}_j^t \right) \right\}_{(i,j) \in J^t} \dot{\cup} \tilde{G}_{\mathcal{O}^1: \mathcal{O}^{t-1}}^{\mathcal{E}} \right);$ // graph
 $I^t \sim \rho_{\tilde{\mathcal{O}} \leftrightarrow \tilde{\mathcal{F}}}^t;$ // for->obj edges
 if $t = \tilde{t}$ **then**
 $I^t := \{(*, \imath)\} \cup (I^t \setminus \{(\cdot, \imath)\});$ // intervention
 for $i = 1$ **to** I **do**
 $\tilde{\mathcal{O}}_i^t := f_O \left(\tilde{\mathcal{O}}_i^{t-1}, \left\{ \tilde{F}_j^t \right\}_{j|(j,i) \in I^t} \right);$ // object nodes
 end
 $\tilde{G}_{\mathcal{O}^1: \mathcal{O}^t} := \left(\tilde{\mathcal{O}}^t \dot{\cup} \tilde{G}_{\mathcal{O}^1: \mathcal{F}^t}^{\mathcal{V}}, \left\{ E \left(\tilde{F}_j^t, \tilde{\mathcal{O}}_i^t \right) \right\}_{(j,i) \in I^t} \dot{\cup} \tilde{G}_{\mathcal{O}^1: \mathcal{F}^t}^{\mathcal{E}} \right);$ // graph
 $J^t \sim \rho_{\tilde{\mathcal{O}}}^t;$ // obj->for edges
 for $j = 1$ **to** J **do**
 $\tilde{F}_j^{t+1} := f_F \left(\tilde{F}_j^t, \left\{ \tilde{\mathcal{O}}_i^t \right\}_{i|(i,j) \in J^t} \right);$ // force nodes
 end
end
return $\tilde{G}_{\mathcal{O}^1: \mathcal{O}^T}$

Appendix C. Model details

C.1. Inductive bias

We introduce two inductive biases: (1) limiting each force node to interact with exactly two objects to reflect pairwise interactions, and (2) enforcing a bidirectional mirroring constraint to ensure temporal coherence in causal attribution. Formally the latter is defined as:

$$\begin{aligned}
 \forall i, j, k, t : \left\{ E(\mathcal{O}_i^t, F_j^{t+1}), E(\mathcal{O}_k^t, F_j^{t+1}) \right\} \subset G_E^t &\implies \\
 E(F_j^{t+1}, \mathcal{O}_i^{t+1}) \in G_E^t \vee E(F_j^{t+1}, \mathcal{O}_k^{t+1}) \in G_E^t, \\
 \forall i, j, t : E(F_j^{t+1}, \mathcal{O}_i^{t+1}) \in G_E^t &\implies E(\mathcal{O}_i^t, F_j^{t+1}) \in G_E^t.
 \end{aligned}$$

C.2. Vector constraints

Causal relevance is coded by C . Perturbing the sub-vectors of the parent with $C = 0$ does not affect the child nodes, i.e., only the causally-relevant $C = 1$ sub-vectors affect the child nodes:

$$\forall t, i : F_j^{t+1, 1PM} = \tilde{F}_j^{t+1, 1PM} \implies f_O \left(O_i^t, \{F_j^{t+1}\}_{j|(j,i) \in I^t}; \theta_O \right) = f_O \left(O_i^t, \{\tilde{F}_j^{t+1}\}_{j|(j,i) \in I^t}; \theta_O \right).$$

Control relevance is coded by P . Two force vectors whose sub-vectors with $P = 1$ are identical have identical control functions that are conditioned on them, i.e., control functions are conditioned only on the control-relevant sub-vectors :

$$\forall t : F_j^{t+1, C1M} = \tilde{F}_j^{t+1, C1M} \implies \rho_{O \leftrightarrow \tilde{F}}^t = \rho_{O \leftrightarrow F}^t.$$

Lastly, *mutability* is coded by M . If $M = 0$, the corresponding sub-vector does not change over time, i.e., an immutable sub-vector does not change over time: $\forall t, j : F_j^{t, CP0} = F_j^{t+1, CP0}$.

C.3. Causal Process Block

Given the data $F^t, O_1^t, \dots, O_n^t$, and the chosen indices J^t , we calculate F^{t+1} in the following way:

$$F^{t+1} := f_F \left(F^t, O_1^t, \dots, O_n^t, J^t \mid \theta_F \right),$$

with O^{t+1} also calculated similarly.

$$\begin{aligned} \text{gate} &:= \frac{1}{|J^t|} \sum_{i \in J^t} \left(O_i^{t, 1PM} W_{\text{gate}}^F \right) W_{\text{output}}^F, \\ \text{residual} &:= \text{gate} + F^{t, CP1}, \\ F^{t+1, CP1} &:= \chi_{\text{gate} \neq 0} \odot \text{FFN}(\text{Norm}(\text{residual})) + \text{residual}, \\ F^{t+1} &:= [F^{t+1, CP1}; F^{t, CP0}], \end{aligned}$$

where $\theta_F := \{W_{\text{gate}}^F, W_{\text{output}}^F, W_1^F, W_2^F, b_1^F, b_2^F\}$ FFN is a feed-forward neural network $\text{FFN}(x) := \max(0, xW_1^F + b_1^F)W_2^F + b_2^F$, W_{gate}^F is the analogue of the attention mechanisms value token projection, and W_{output}^F is again the analogous out-projection that maps the token from the attention dimension back to residual dimension (Vaswani et al., 2017)

Appendix D. Optimization details

During controller optimization, the implementation uses several auxiliary regularizers and stabilizers that are not part of the main algorithmic description. An edge-activation cost penalizes every edge that is actually executed, encouraging sparse graphs. A blocked-edge penalty penalizes an edge that was proposed by the controller but later rejected by legality checks or post-processing before execution. An effect-size threshold suppresses any proposed edge whose learned causal-force vector has norm below a fixed threshold, treating it as too weak to contribute to state propagation; an effect-size penalty then penalizes selecting such weak edges in the first place. The controller objective also includes an entropy bonus, which encourages exploration by discouraging premature collapse to

deterministic policies. During training, both controllers share a temperature schedule that anneals linearly from 2.0 to 1.0, whereas during validation and evaluation the temperature is fixed to 1.0 and controller decisions are deterministic by default. Temporal-difference learning further uses slowly updated target controllers via Polyak averaging (Polyak and Juditsky, 1992), and stored Q-values are clipped to the range $[-10, 10]$ for stability. In the default configuration, the edge-activation cost is 0.1, the blocked-edge penalty is 0.1, the effect-size threshold is 0.1, the effect-size penalty is 0.15, and the entropy-bonus coefficient is 0.1. Importantly, an edge that is proposed but later blocked is still kept in training as a negative example rather than being discarded.

A further implementation detail concerns how sparse acyclic interaction graphs are enforced during recurrent rollout. The modeled scene entities and the pairwise forces between them are referred to as endogenous objects and endogenous forces. Separate exogenous objects and exogenous forces represent influences originating outside that modeled object set. In the default configuration there is a single exogenous object, intended to represent the surrounding environment or boundary effects, such as walls. Graph construction proceeds over recurrent steps using a two-tier frontier, where a frontier is the set of object nodes currently eligible to participate in the next wave of interaction propagation. The endogenous frontier is defined by sink nodes of currently executed endogenous edges, where a sink is the object receiving the effect of a directed interaction. If no such sink exists at the beginning of rollout, the object directly targeted by the action is used as the initial seed frontier. Endogenous edge proposals are restricted to object pairs touching the current endogenous frontier, and their directions are further constrained so that frontier nodes act as sources for subsequent diffusion rather than immediate sinks. After an object is directly intervened on by the action, incoming forces to that object are suppressed for the next recurrent step, so the intervened object temporarily seeds propagation without itself being immediately re-entered as a sink. The exogenous frontier is cumulative: it contains all endogenous objects that have ever been intervention targets or sinks in earlier recurrent steps. Exogenous edges may target only this cumulative set, and only when the endogenous frontier is empty or no endogenous edge is active. Directed-acyclic-graph (DAG) structure is enforced at two levels: first, candidate endogenous edges whose admissible directions would introduce a cycle are masked before activation; second, after direction selection, any direction that would violate acyclicity is flipped to the valid alternative when one exists, and otherwise the edge is deactivated. A final sequential reachability update ensures that the executed directed graph remains acyclic throughout the rollout.

Appendix E. Plots

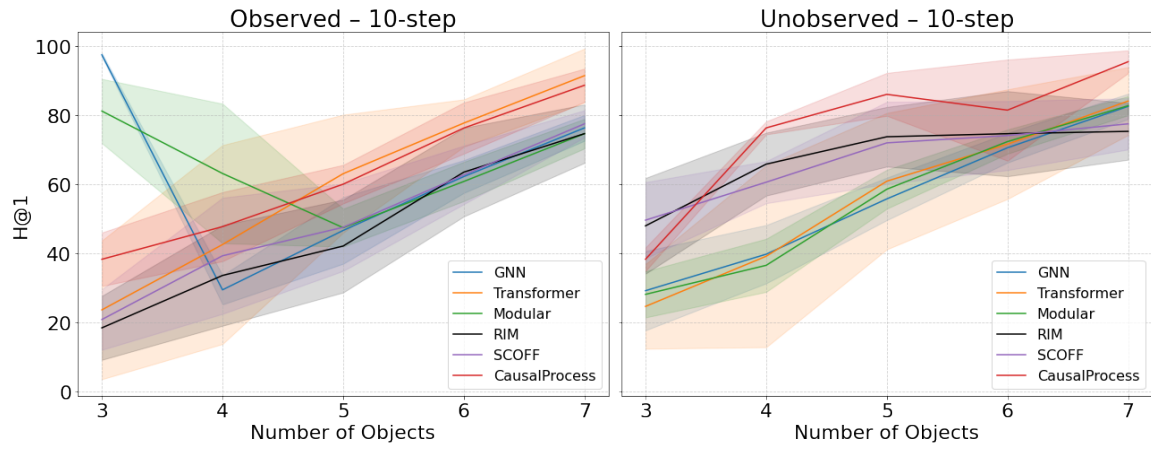


Figure 7: Prediction metric vs number of objects for 10-steps.

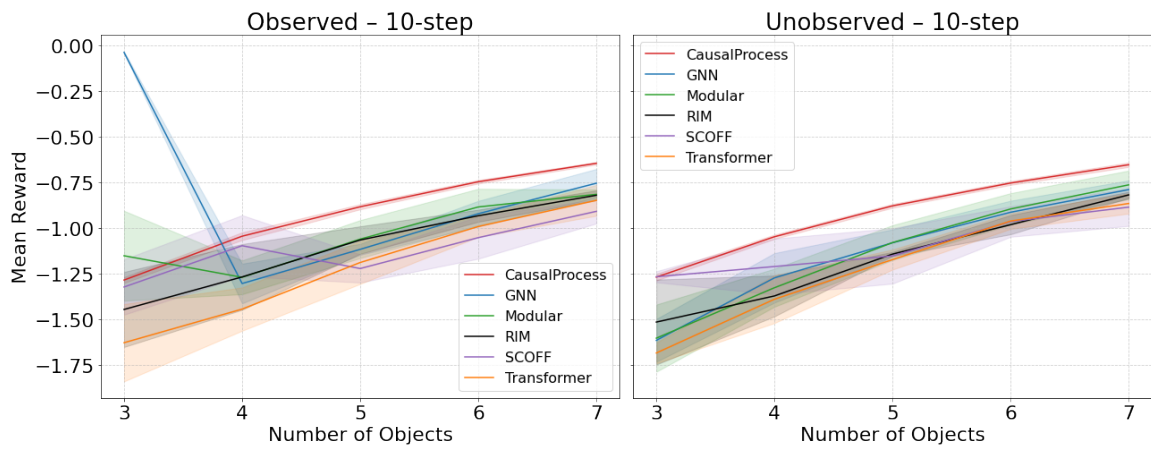


Figure 8: *Downstream RL results over number of objec.* Mean reward vs number of objects. All results are the average of 10 seeds

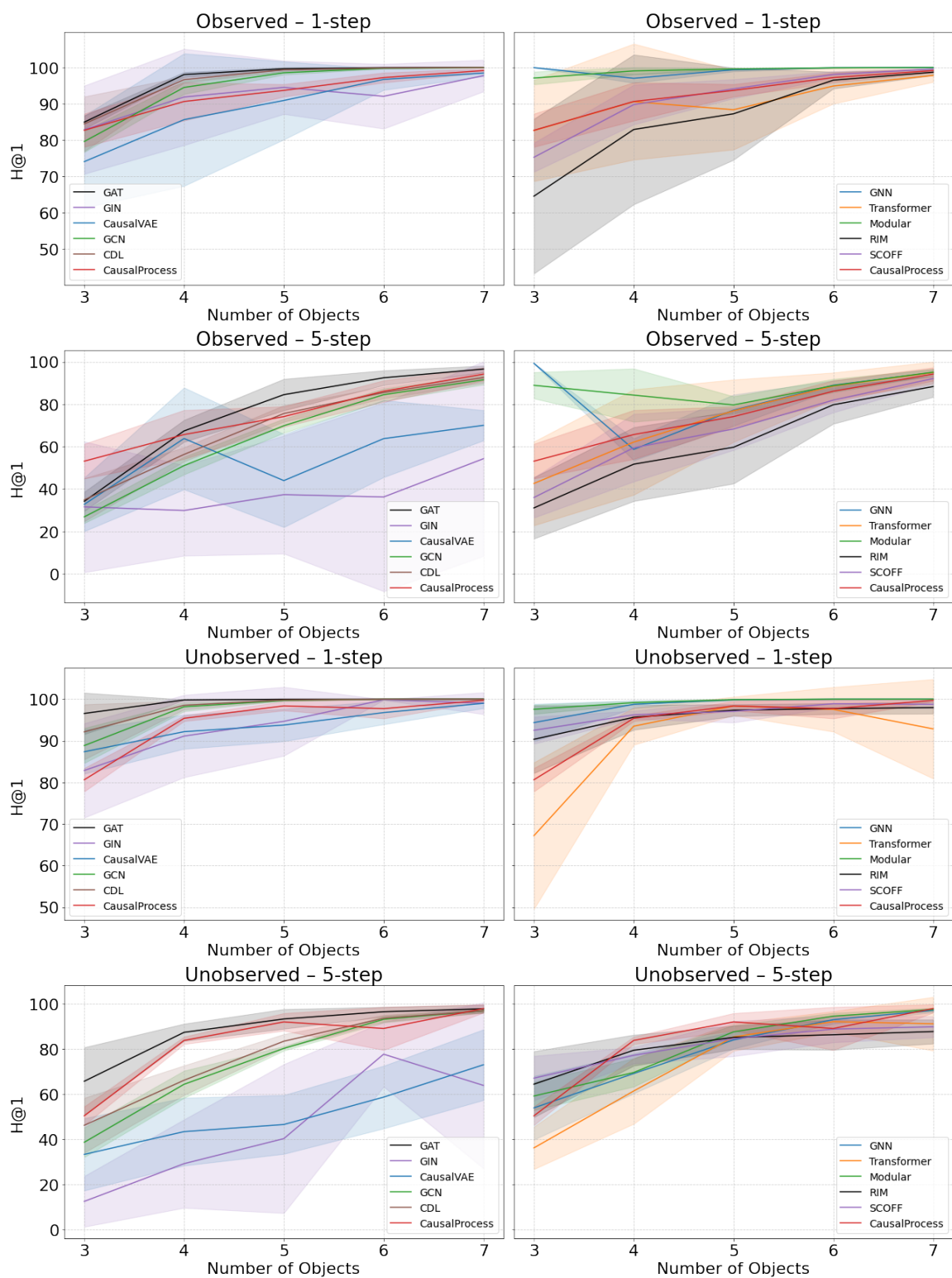


Figure 9: Prediction metric vs number of objects after 1 and 5 steps (average of 10 seeds).

CAUSAL PROCESS MODELS

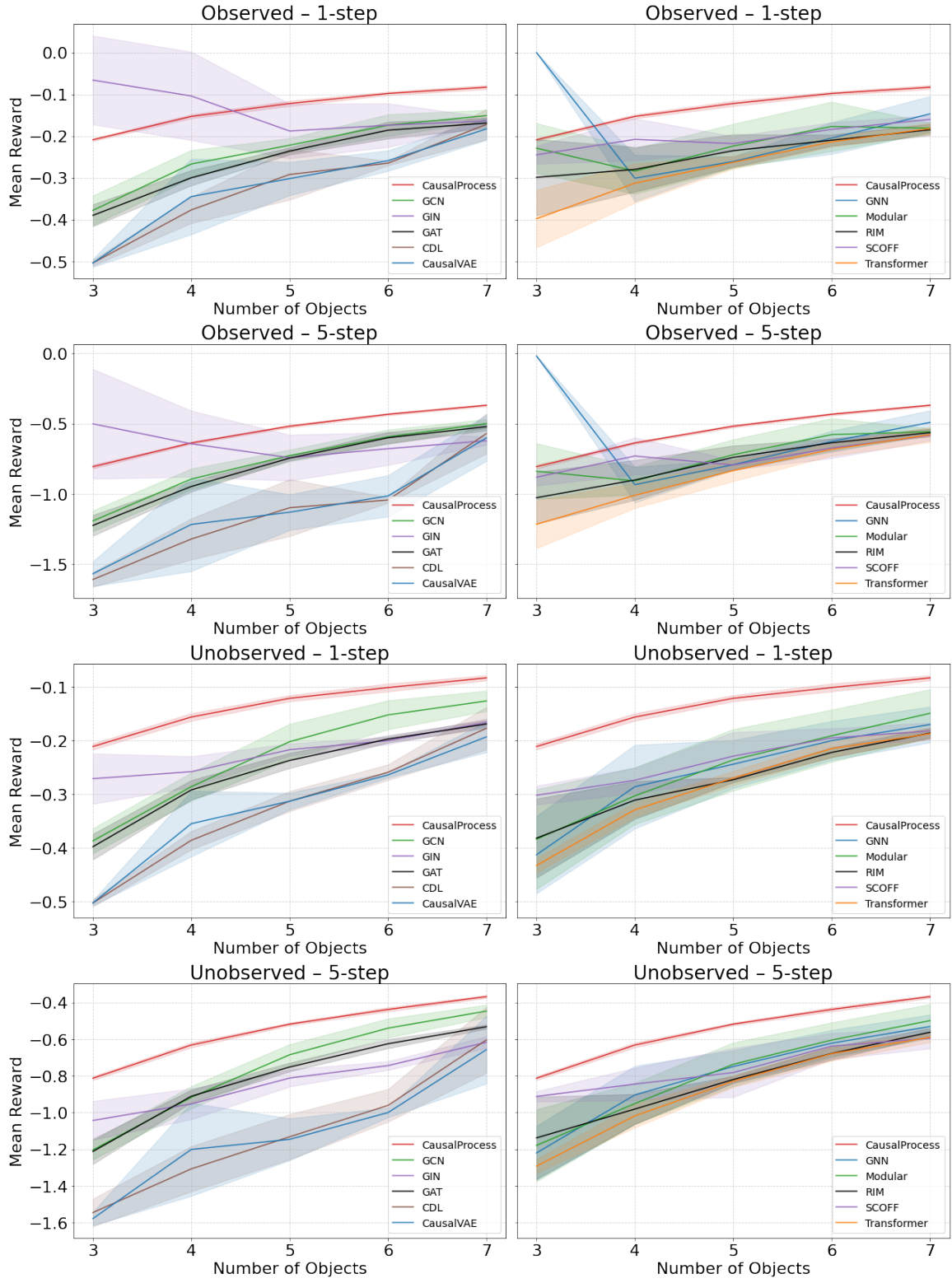


Figure 10: Mean reward vs number of objects after 1 and 5 steps (average of 10 seeds).

Appendix F. Ablation Study

The CPM combines (i) a structured latent factorization, and (ii) an RL-based procedure for selecting sparse, time-varying causal edges. To isolate which components are essential for performance, we conduct targeted ablations of the three most central design choices: learned reward shaping, learned causal controllers, and the structured C/P/M representation.

Setup. All ablations are evaluated in the **Observed** setting with **3 objects**, averaged over **2 random seeds**. We report both (i) prediction accuracy via Hits@1 (H@1; Sec. 5.2) and (ii) downstream model-based RL performance via mean environment reward (Sec. 5.4), across rollout horizons $K \in \{1, 5, 10\}$. We report *deltas* computed as $\Delta = (\text{ablated}) - (\text{full CPM})$, so negative values indicate degradation relative to the full model. For each ablation, we keep the training pipeline and hyperparameters fixed and disable exactly one component, leaving all others intact.

Reward ablated (zero reward shaping). We replace the learned reward models R_O and $R_{O \leftrightarrow F}$ (Eq. 4) with constant zero outputs. This removes the inverse-RL reward shaping signal that otherwise trains the controllers to prefer graph decisions that reduce prediction error.

Controller ablated (random graph decisions). We replace the learned controller policies π_O and $\pi_{O \leftrightarrow F}$ (Sec. 4.2) with uniform random edge selection (while keeping the same structural constraints, e.g., pairwise interactions / mirroring when applicable). This removes the model’s ability to adaptively select sparse causal edges based on the current state.

C/P/M regions ablated (fused latent). We remove the structured factorization into Causal relevance (C), Control relevance (P), and Mutability (M) (Sec. 4.1.1) and instead use a single undifferentiated latent vector with the same total dimensionality. This ablates the hard architectural routing bias that forces semantically distinct subspaces.

Results and discussion. Table 1 summarizes the results. All ablations consistently degrade both prediction and downstream planning, and the effect *amplifies strongly with rollout horizon*. At $K = 1$, the H@1 drops are small (≤ 0.27 percentage points), while at $K = 10$ they become substantial (up to -13.01 H@1 points and -1.08 reward). This pattern suggests that the components are particularly important in the long-horizon regime where compounding error is the dominant failure mode.

Among the components, the structured **C/P/M factorization** is most critical at long horizons ($\Delta\text{H@1} = -13.01$ at $K = 10$; $\Delta\text{Reward} = -1.08$), consistent with its intended role of separating immutable vs. mutable and causally relevant vs. irrelevant factors during multi-step rollouts. The **reward shaping** ablation yields the second-largest long-horizon degradation in H@1 (-10.63 at $K = 10$), supporting the claim that learned rewards are important for training the controllers toward causally useful graphs. The **random-controller** ablation shows smaller but still meaningful degradation (-8.45 H@1 at $K = 10$), which we note may be partially understated in the 3-object setting because the space of possible object pairs is small; with more objects, random edge selection becomes less likely to pick the correct interacting pair.

Table 1: **Ablation results (Observed, 3 objects, mean over 2 seeds).** Δ values are computed as *ablated* – *full CPM*. Negative indicates worse.

Ablation	Δ H@1 (% points)			Δ Reward		
	1-step	5-step	10-step	1-step	5-step	10-step
Reward ablated	−0.27	−3.97	−10.63	−0.20	−0.61	−0.74
Controller ablated	−0.17	−3.52	−8.45	−0.09	−0.21	−0.24
C/P/M ablated	−0.19	−4.12	−13.01	−0.24	−0.85	−1.08

Appendix G. Graph Recovery and Structured Representation Analysis (Observed 3-object setting)

G.1. Linear probing of the structured (C, P, M) subspaces

We test whether *object location* is preferentially encoded in the *mutable* ($M=1$) subspaces, consistent with the intended semantics of M (Sec. 4.1.1): positions change over time, whereas immutable properties should not.

Protocol. For each trained CPM (two seeds), we freeze all model parameters and extract the object latents O_i^t , which are partitioned into the eight sub-vectors corresponding to all $(C, P, M) \in \{0, 1\}^3$:

$$O_i^t = \bigoplus_{C,P,M \in \{0,1\}} O_i^{t,CPM}.$$

For each sub-vector (e.g., $O_i^{t,CPM}$, $O_i^{t,CM}$, $\dots$, $O_i^{t,\emptyset}$), we train a linear regressor to predict object location and report held-out test R^2 . We use the same train/test split across seeds and report mean $\pm$ std over seeds.

Result. Location is almost perfectly decodable from all subspaces with $M=1$ (CPM/CM/PM/M), with $R^2 \approx 0.99$, while decodability collapses in $M=0$ subspaces (CP/C/P/ $\emptyset$), with $R^2 \approx 0.08\text{--}0.14$ on average (Table 2). This supports that the model routes positional information into the **mutable** channel, aligning with the inductive bias described in Sec. 4.1.1.

G.2. Graph recovery against the ground-truth interaction graph

To provide a quantitative measure of causal graph discovery, we compare the learned directed interaction graph to the simulator’s ground-truth interaction graph in the observed 3-object setting.

Ground truth. In the [Ke et al. \(2021\)](#) environment used in our experiments, interactions are sparse and directed: the only force is pushing (double-pushes are disallowed), and only heavier objects can push lighter ones. Thus, when an interaction occurs, the ground-truth directed edge is always *heavy*→*light*. (In the observed setting, weight corresponds to color intensity.)

Predicted graph and metrics. At each timestep, CPM predicts a sparse directed interaction graph. We threshold edge activations at 0.5 to obtain a discrete predicted graph and compute: (i) Precision/Recall/F1 for edge presence vs. ground truth, and (ii) **Direction validity** = fraction of predicted edges that follow the correct heavy→light direction. We report metrics both over *all*

Region	Location R^2 (mean $\pm$ std)
M	0.991 ± 0.003
PM	0.991 ± 0.003
CM	0.991 ± 0.003
CPM	0.990 ± 0.003
$\emptyset$	0.141 ± 0.137
C	0.115 ± 0.087
CP	0.091 ± 0.088
P	0.078 ± 0.084

Table 2: Linear probing of CPM object subspaces for **location**. We train a linear regressor to predict object position from each (C, P, M) sub-vector. The combined panel reports mean $\pm$ std R^2 over two seeds. Location is almost perfectly decodable from all subspaces with $M=1$ (CPM/CM/PM/M), but not from $M=0$ subspaces (CP/C/P/ $\emptyset$), indicating that positional information is routed into the **mutable** channel. Values are test R^2 (mean $\pm$ std over two seeds).

Evaluation set	Precision	Recall	F1	Direction validity
All steps	0.084 ± 0.010	0.403 ± 0.059	0.139 ± 0.018	0.738 ± 0.099
Interaction steps only	0.544 ± 0.050	0.403 ± 0.059	0.463 ± 0.057	0.853 ± 0.088

Table 3: Graph recovery in the observed 3-object setting (mean $\pm$ std over two seeds). Direction validity measures the fraction of predicted edges that follow the correct heavy $\rightarrow$ light direction.

steps (including non-interaction timesteps) and over *interaction steps only* (timesteps where the ground-truth graph is non-empty).

Result. On interaction timesteps, CPM recovers directed interaction edges with substantially higher precision and predicts the correct heavy $\rightarrow$ light direction in the large majority of predicted edges (Table 3). Over all timesteps, precision is lower due to false positives on non-interaction steps, suggesting that sparsity calibration is an important direction for future improvement.

G.3. Visualization of the learned directed graph

To provide an interpretable summary of the learned graph structure, we average the model’s predicted edge probabilities $p(i \rightarrow j)$ over the evaluation dataset. For each unordered pair $\{i, j\}$, we draw the arrow in the more probable direction and label it with the *directional share*:

$$\frac{p(i \rightarrow j)}{p(i \rightarrow j) + p(j \rightarrow i)} \times 100\%.$$

Fig. 14 shows per-seed graphs and their average. Nodes are colored by object weight (darker indicates heavier), so heavy $\rightarrow$ light corresponds to dark $\rightarrow$ light.

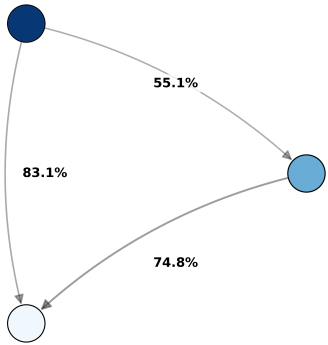
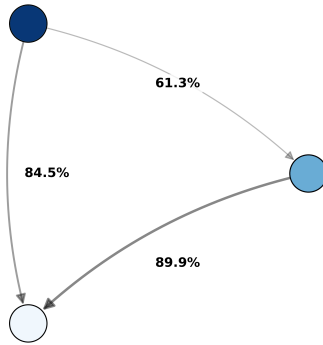
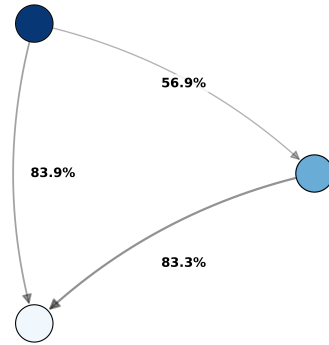
Figure 11: 1st Learned DAG.Figure 12: 2nd Learned DAG.

Figure 13: Average

Figure 14: Learned directed interaction graph in the observed 3-object setting. For each pair of objects $\{i, j\}$ we average predicted edge probabilities $p(i \rightarrow j)$ over the dataset, draw the arrow in the more probable direction, and label it with the directional share $\frac{p(i \rightarrow j)}{p(i \rightarrow j) + p(j \rightarrow i)}$. This visualization qualitatively matches the quantitative directionality reported in Table 3.