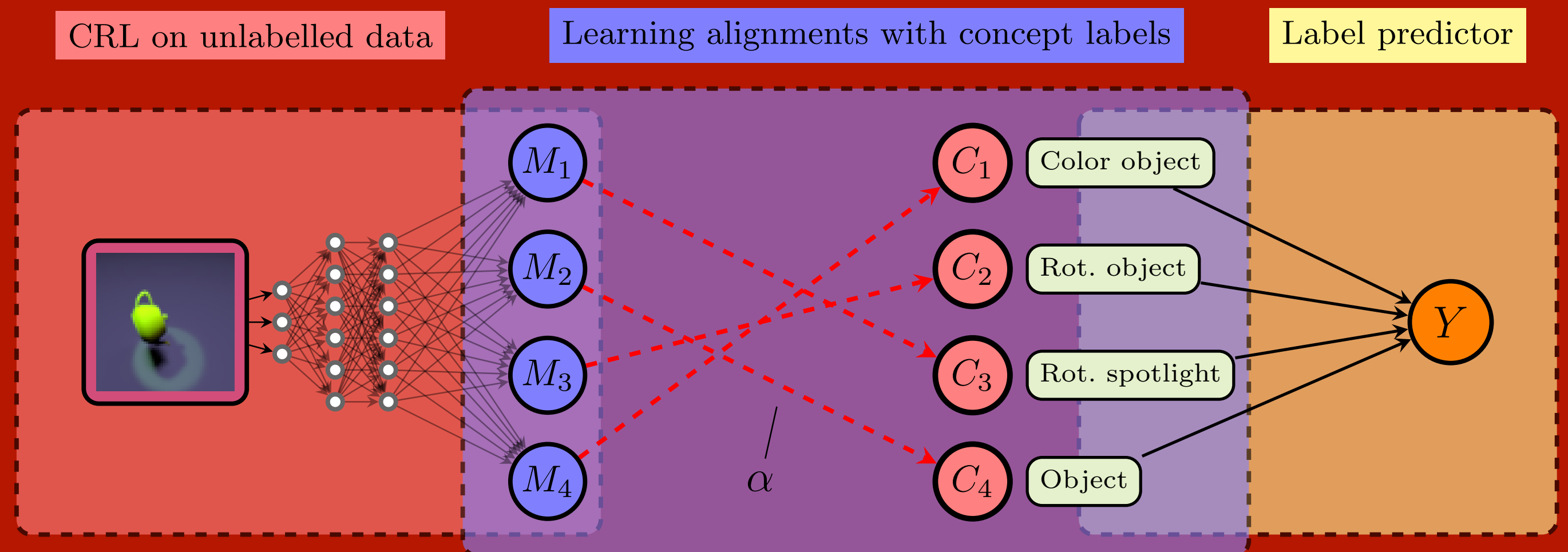


Sample-efficient Learning of Concepts with Theoretical Guarantees:

from Data to Concepts without Interventions



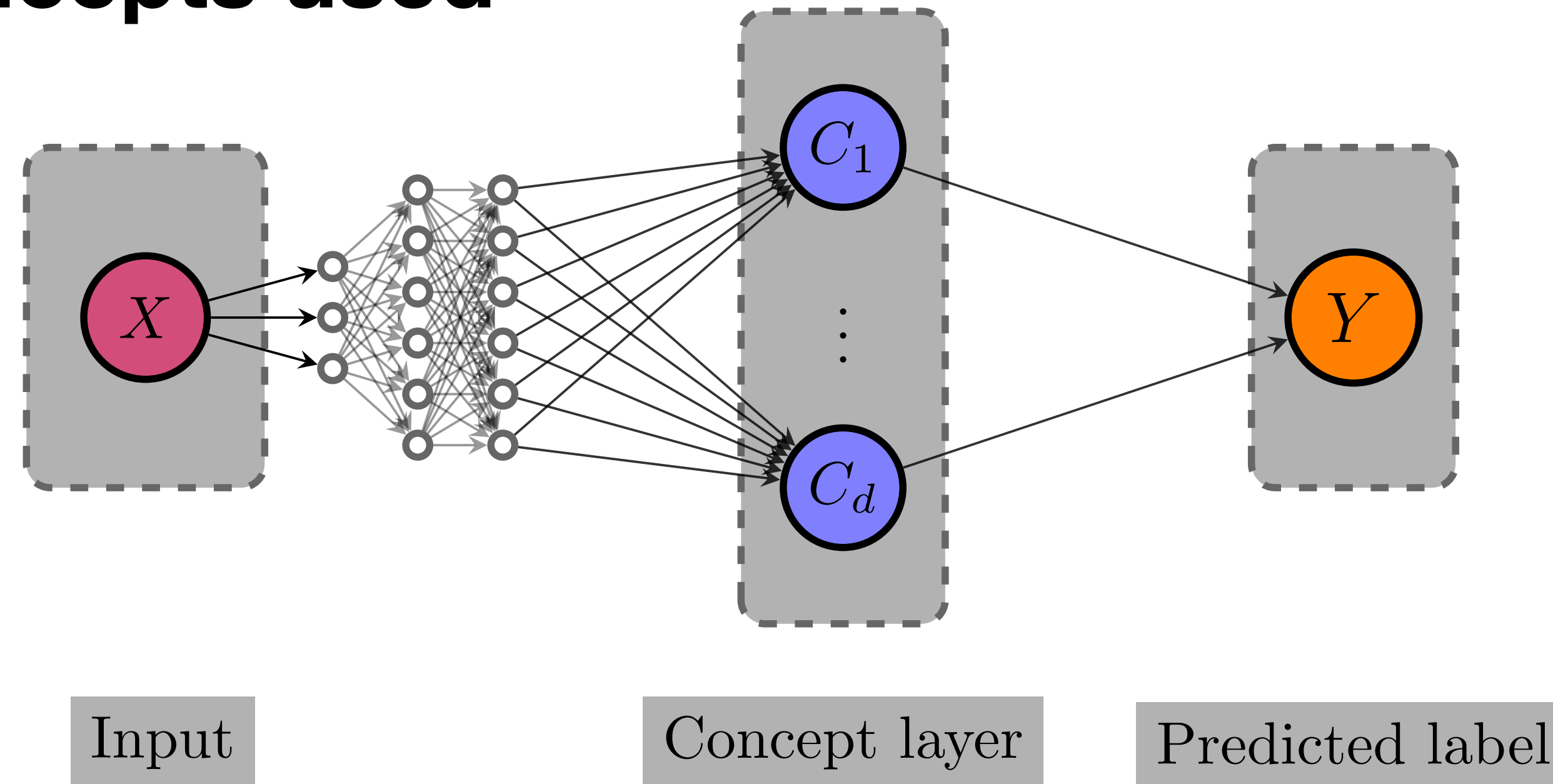
Outline

- ▶ **Concept-based Models**
- ▶ **Causal Representation Learning**
- ▶ **Combining Concepts and Causal Variables**
- ▶ **Alignment learning**
- ▶ **Experimental Results**

Concept-Based Models

XAI and the need for high-level concepts

How are concepts used



► An explanation is given as:

$$\{ (h(x)_c, \psi_{cj}) \mid c \in \{1, \dots, d\} \}$$

- c is the concept index
- ψ linear activations
- $h(x)_c$ represents the strengt of concept c in the input

Concept-Based Examples

Concept-Based Models

- ▶ Concept Bottleneck Models, [Koh et al., 2018]
- ▶ Concept Embedding Models [Zarlenga et al., 2022]
- ▶ GLANCE-Nets [Marconato et al, 2022]

Concept-Based Explanations

- ▶ (Non)-linear probing [Alain & Bengio, 2016]
- ▶ TCAV [Kim et al, 2018]
- ▶ Sparse Auto-Encoders [Sharkey & Milidge, 2022]
- ▶ Mechanistic interpretability

Causal Representation Learning

Causal Representation Learning

Brief introduction

To overcome this, we can assume a **causal setting**:

- ▶ There are 2 variables $G_1 \in \{\text{Cow}, \dots\}$, $G_2 \in \{\text{Mountains, beach}, \dots\}$
- ▶ Nature provides some function $f(G_1, G_2) = X \in \mathbb{R}^{h \times w}$ that generates the image

CRL methods try to learn the **reverse process**, given a dataset of only images

Causal Representation Learning

Brief introduction

To overcome this, we can assume a **causal setting**:

- ▶ There are 2 variables $G_1 \in \{\text{Cow}, \dots\}$, $G_2 \in \{\text{Mountains, beach}, \dots\}$
- ▶ Nature provides some function $f(G_1, G_2) = X \in \mathbb{R}^{h \times w}$ that generates the image

CRL methods try to learn the **reverse process**, given a dataset of only images

- ▶ Learn an encoder $g_\phi(X) = M \approx G$
- ▶ Typically a Variational Auto Encoder (VAE)

Causal Representation Learning

Caveats

- ▶ Many assumptions needed
- ▶ Even then, in general this problem is **non-identifiable**
- ▶ It is identifiable up to an equivalence class ($\sim$)
- ▶ The learned M are:
 - **Permuted**,
 - **Component-wise transformations** of G .

$$PT(M) = \begin{bmatrix} T_{\pi(1)}(M_{\pi(1)}) \\ T_{\pi(2)}(M_{\pi(2)}) \\ \vdots \\ T_{\pi(d)}(M_{\pi(d)}) \end{bmatrix} = \begin{bmatrix} G_1 \\ G_2 \\ \vdots \\ G_d \end{bmatrix} = G$$

Combining Concepts and Causal Variables

Combining Concepts and Causal variables

Main idea

We will assume that the **causal variable** and **interpretable concepts** are the same!

Problem: The learned causal variables are **permuted** and **transformed**

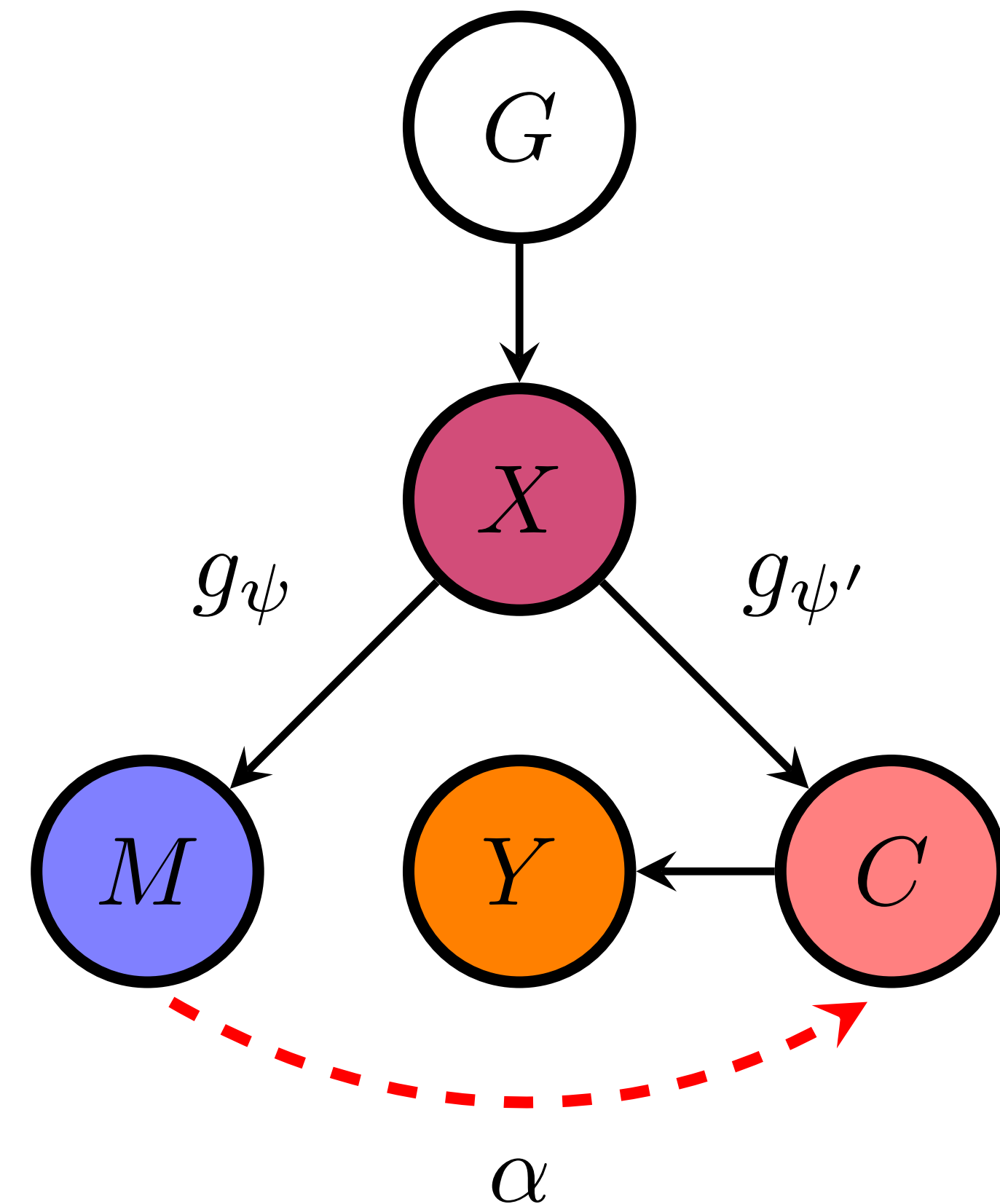
Identifiability theory

Combined with concepts

Learning $\alpha: M \mapsto C$

$$PT(M) = \begin{bmatrix} T_{\pi(1)}(M_{\pi(1)}) \\ T_{\pi(1)}(M_{\pi(2)}) \\ \vdots \\ T_{\pi(d)}(M_{\pi(d)}) \end{bmatrix} = \begin{bmatrix} C_1 \\ C_2 \\ \vdots \\ C_d \end{bmatrix} = C$$

Involves learning a **permutation** and a **1-dimensional mapping**



Alignment Learning

Alignment Learning

Intuition

For now imagine 2 variables and 2 outcomes:

$$\begin{array}{l} Y_1 = \beta_2 X_2 \\ Y_2 = \beta_1 X_1 \end{array} \quad \longrightarrow \quad \begin{bmatrix} Y_1 \\ Y_2 \end{bmatrix} = \begin{bmatrix} 0 & \beta_2 \\ \beta_1 & 0 \end{bmatrix} \begin{bmatrix} X_1 \\ X_2 \end{bmatrix}$$

Alignment Learning

Intuition

For now imagine 2 variables and 2 outcomes:

$$\begin{array}{l} Y_1 = \beta_2 X_2 \\ Y_2 = \beta_1 X_1 \end{array} \quad \longrightarrow \quad \begin{bmatrix} Y_1 \\ Y_2 \end{bmatrix} = \begin{bmatrix} 0 & \beta_2 \\ \beta_1 & 0 \end{bmatrix} \begin{bmatrix} X_1 \\ X_2 \end{bmatrix}$$

Regressing $X = (X_1, X_2)$ simultaneously on to $Y = (Y_1, Y_2)$ would give

$$\begin{bmatrix} Y_1 \\ Y_2 \end{bmatrix} = \begin{bmatrix} \hat{\beta}_1^1 & \hat{\beta}_2^1 \\ \hat{\beta}_1^2 & \hat{\beta}_2^2 \end{bmatrix} \begin{bmatrix} X_1 \\ X_2 \end{bmatrix}, \text{ with } \hat{\beta}_1^1 \approx 0, \hat{\beta}_2^2 \approx 0$$

New Concept-Based Framework

Thank you for your attention!

