

Nonparametric Partial Disentanglement via Mechanism Sparsity: Sparse Actions, Interventions and Sparse Temporal Dependencies

Sébastien Lachapelle

Samsung AI Lab, Montreal; Mila & DIRO, Université de Montréal

S.LACHAPELLE@SAMSUNG.COM

Pau Rodríguez López

ServiceNow Research

Yash Sharma

Tübingen AI Center, University of Tübingen

Katie Everett

MIT CSAIL

Rémi Le Priol

Mila & DIRO, Université de Montréal

Alexandre Lacoste

ServiceNow Research

Simon Lacoste-Julien

Samsung AI Lab, Montreal; Mila & DIRO, Université de Montréal; Canada CIFAR AI Chair

Editor: Elias Bareinboim

Abstract

This work introduces a novel principle for disentanglement we call *mechanism sparsity regularization*, which applies when the latent factors of interest depend sparsely on observed auxiliary variables and/or past latent factors. We propose a representation learning method that induces disentanglement by *simultaneously* learning the latent factors and the sparse causal graphical model that explains them. We develop a nonparametric identifiability theory that formalizes this principle and shows that the latent factors can be recovered by regularizing the learned causal graph to be sparse, under some assumptions such as the absence of instantaneous causal effects between latent factors. More precisely, we show identifiability up to a novel equivalence relation we call *consistency*, which allows some latent factors to remain entangled (hence the term *partial* disentanglement). To describe the structure of this entanglement, we introduce the notions of *entanglement graphs* and *graph preserving functions*. We further provide a graphical criterion which guarantees *complete* disentanglement, that is identifiability up to permutations and element-wise transformations. We demonstrate the scope of the mechanism sparsity principle as well as the assumptions it relies on with several worked out examples. For instance, the framework shows how one can leverage multi-node interventions with unknown targets on the latent factors to disentangle them. We further draw connections between our nonparametric results and the now popular exponential family assumption. Lastly, we propose an estimation procedure based on variational autoencoders and a sparsity constraint and demonstrate it on various synthetic datasets. This work is meant to be a significantly extended version of [Lachapelle et al. \(2022\)](#).

Keywords: identifiable representation learning, causal representation learning, disentanglement, nonlinear independent component analysis, causal discovery

Contents

1	Introduction	4
2	Problem Setting, Entanglement Graphs & Disentanglement	7
2.1	An Identifiable Latent Causal Model	7
2.2	Entanglement Maps & Entanglement Graphs	10
2.3	Identifiability and Observational Equivalence	11
2.4	Equivalence up to Diffeomorphism	11
2.5	Disentanglement and Equivalence up to Permutation	13
3	Nonparametric Partial Disentanglement via Mechanism Sparsity	14
3.1	Conditional Independence is Insufficient for Disentanglement and Graph Discovery	14
3.2	A First Mathematical Insight for Disentanglement via Mechanism Sparsity	15
3.3	Graph Preserving Maps	17
3.4	Nonparametric Identifiability via Auxiliary Variables with Sparse Influence	19
3.4.1	Unknown-Target Interventions on the Latent Factors	23
3.5	Nonparametric Identifiability via Sparse Temporal Dependencies	24
3.6	Combining Sparsity Regularization on $\hat{G}^a$ & $\hat{G}^z$	26
3.7	Graphical Criterion for Complete Disentanglement	26
3.8	Proofs of Theorems 1, 2 & 3 and their Sufficient Influence Assumptions	27
3.8.1	Sufficient influence assumption of Theorem 1 and its proof	28
3.8.2	Sufficient Influence Assumption of Theorem 2 and its Proof	29
3.8.3	Sufficient Influence Assumption of Theorem 3 and its Proof	31
3.9	Examples to Illustrate the Scope of the Theory	32
3.9.1	Continuous Auxiliary Variables (Theorem 1)	33
3.9.2	Discrete Auxiliary Variables or Interventions (Theorem 2)	34
3.9.3	Temporal Dependencies (Theorem 3)	35
4	Partial Disentanglement via Mechanism Sparsity in Exponential Families	37
4.1	Exponential Family Latent Transition Models	37
4.2	Conditions for Quasi-Linear Identifiability	37
4.3	Partial Disentanglement via Sparse Time Dependencies in Exponential Families	39
5	Model Estimation with a Sparsity Constraint	41
6	Evaluation with R_{con} and SHD	43
7	Related Work	44
8	Experiments	47
8.1	Graphs Allowing Complete Disentanglement (Satisfying Assumption 5)	48
8.2	Graphs Allowing only Partial Disentanglement (Violating Assumption 5)	51
9	Conclusion	54

A	Identifiability Theory - Nonparametric Case	57
A.1	Useful Lemmas	57
A.2	Proof of Proposition 1	58
A.3	Proof of Proposition 2	58
A.4	The Consistency Relations (Definitions 13 & 14) are Equivalence Relations	61
A.4.1	Combining Equivalence Relations	63
A.5	Technical Lemmas Leading to Theorems 1, 2 & 3	64
A.6	Connecting to the Graphical Criterion of Lachapelle et al. (2022)	67
B	Identifiability Theory - Exponential Family Case	68
B.1	Technical Lemmas and Definitions	68
B.2	Proof of Linear Identifiability (Theorem 4)	69
B.3	Proof of Theorem 5	71
B.4	Connecting to Sufficient Influence Assumptions of Lachapelle et al. (2022)	73
C	Experiments	74
C.1	Synthetic Datasets	74
C.1.1	Datasets Satisfying the Graphical Criterion	75
C.1.2	Datasets that Violate the Graphical Criterion	77
C.2	Implementation Details of our Constrained VAE Method	78
C.3	Baselines	78
C.4	Unsupervised Hyperparameter Selection	79
D	Miscellaneous	80
D.1	On the Invertibility of the Mixing Function	80
D.2	Contrasting with the Assumptions of Khemakhem et al. (2020a) & Yao et al. (2022b)	81
D.3	Derivation of the ELBO	82
E	Author Contributions	83
E.1	Contributions to the Extended Version	83
E.2	Contributions to the CLear Version (Lachapelle et al., 2022)	83

1 Introduction

It has been proposed that causal reasoning will be central in moving modern machine learning algorithms beyond their current shortcomings, such as their lack of *robustness*, *transferability*, and *interpretability* (Pearl, 2019; Schölkopf, 2019; Goyal and Bengio, 2021). To achieve this, the field of *causal representation learning* (CRL) (Schölkopf et al., 2021) aims to learn representations of high-dimensional observations, such as images, that are suitable to perform causal reasoning, such as predicting the effect of unseen interventions and answering counterfactual queries. A popular formalism to do so is to assume that the observations $\mathbf{x} \in \mathbb{R}^{d_x}$ are sampled from a generative model of the form $\mathbf{x} = \mathbf{f}(\mathbf{z})$ where $\mathbf{z} \in \mathbb{R}^{d_z}$ is a random vector of *unobserved* and *semantically meaningful* variables, also called latent factors, distributed according to an unknown *causal graphical model* (CGM) (Pearl, 2009; Peters et al., 2017) and transformed by a potentially highly nonlinear *decoder*, or *mixing function*, $\mathbf{f}$ (Kocaoglu et al., 2018; Volodin, 2021; Lachapelle et al., 2022; Lippe et al., 2023b; Brehmer et al., 2022; Ahuja et al., 2023; Buchholz et al., 2023; von Kügelgen et al., 2023; Zhang et al., 2023; Jiang and Aragam, 2023). The goal is then to recover the latent factors $\mathbf{z}_i$ up to permutation and rescaling, as well as the causal relationships that explain them. This is closely related to the problem of *disentanglement* (Bengio et al., 2013; Higgins et al., 2017; Locatello et al., 2020), which also aims to extract interpretable variables from high-dimensional observations, but without the emphasis on modeling their causal relations. Such problems are plagued by the difficult question of *identifiability*, which is of crucial importance in the classical settings of *causal discovery* (Pearl, 2009; Peters et al., 2017), where $\mathbf{f}$ is assumed to be the identity, and *independent component analysis* (ICA) (Hyvärinen et al., 2001, 2023), where the causal graph over latents is assumed empty. In the former, one can only identify the Markov equivalence class of the causal graph (assuming faithfulness), thus leaving some edge orientations ambiguous (Pearl, 2009), while in the latter, identifiability of the ground-truth latent factors is impossible when assuming a general nonlinear $\mathbf{f}$, (Hyvärinen and Pajunen, 1999). The general CRL problem inherits the difficulties from both of these settings, which makes identifiability especially challenging. Various strategies to improve identifiability have been contributed to the literature, such as assuming access to *interventional data* in which latent factors are targeted by interventions (Lachapelle et al., 2022; Lippe et al., 2022, 2023b; Ahuja et al., 2023), or access to an *auxiliary variable* $\mathbf{a}$ that makes factors $\mathbf{z}_i$ mutually independent when conditioned on (Hyvärinen et al., 2019; Khemakhem et al., 2020a,b). A valid auxiliary variable $\mathbf{a}$ must be observed and could correspond, for instance, to a time or environment index, an action in an interactive environment, or even a previous observation if the data have temporal structure. See Section 7 for a more extensive review of existing approaches to the identification of latent variables.

The present paper introduces¹ *mechanism sparsity regularization* as a new principle for the identification of latent variables. We show that if (i) an auxiliary variable $\mathbf{a}$ is observed and affects the latent variables *sparsely* and/or (ii) the latent variables $\mathbf{z}$ present *sparse* temporal dependencies, then the latent variables can be recovered by learning a graphical model for $\mathbf{z}$ and $\mathbf{a}$ and regularizing it to be sparse (Theorems 1, 2, 3 & 5). More specifically, we consider models of the form $\mathbf{x}^t = \mathbf{f}(\mathbf{z}^t) + \mathbf{n}^t$, where $\mathbf{n}^t$ is independent noise (Assumption 1) and the latent factors $\mathbf{z}_i^t$ are mutually independent given the past factors and the auxiliary variables, i.e. $p(\mathbf{z}^t | \mathbf{z}^{<t}, \mathbf{a}^{<t}) = \prod_{i=1}^{d_z} p(\mathbf{z}_i^t | \mathbf{z}^{<t}, \mathbf{a}^{<t})$ (Assumption 2). Crucially, we leverage the assumption that these mechanisms are sparse in the sense that $p(\mathbf{z}^t | \mathbf{z}^{<t}, \mathbf{a}^{<t})$ factorizes according to a sparse

1. A shorter version of this work originally appeared in Lachapelle et al. (2022).

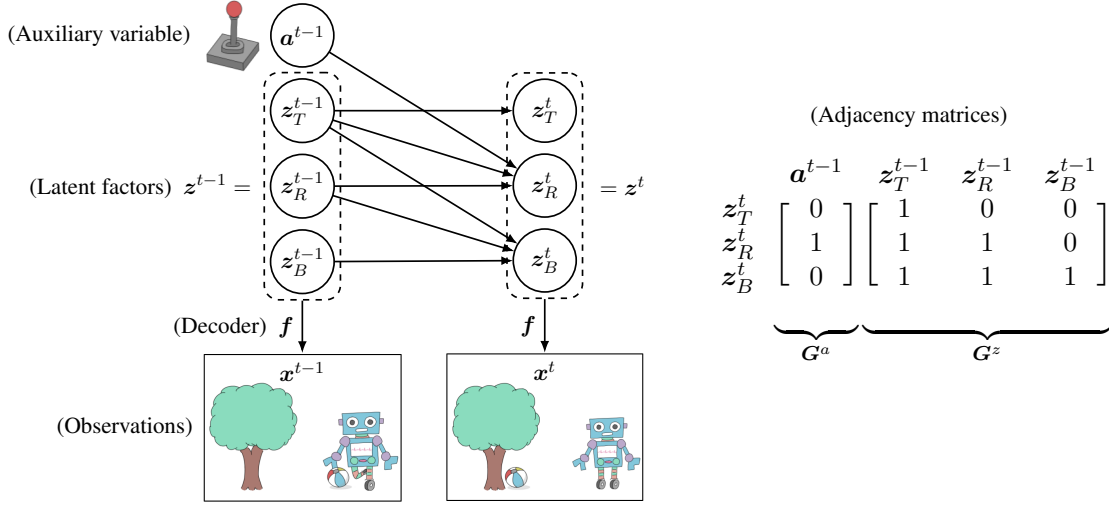


Figure 1: A minimal motivating example. The latent factors z_T^t , z_R^t and z_B^t represent the x -positions of the tree, the robot and the ball at time t , respectively. Only the image of the scene x^t and the action a^{t-1} are observed. See end of Section 2.1 for details.

causal graph G (Assumption 3). Interestingly, if a corresponds to an intervention index, our framework explains how interventions targeting unknown subsets of latent factors can identify them (Section 3.4.1). We emphasize that the settings where the data have no temporal dependencies or no auxiliary variable a are special cases of our framework. Our identifiability results are summarized in Table 1.

This work is meant to be an extended version of Lachapelle et al. (2022) in which we generalize along two main axes: First, we relax the *exponential family* assumption by providing a fully *nonparametric* treatment. Secondly, our results drop the graphical criterion of Lachapelle et al. (2022) and, thus, allow for *arbitrary* latent causal graphs. As a consequence of this relaxation, instead of guaranteeing identifiability up to permutation and element-wise transformation, we guarantee identifiability up to what we call *a-consistency* or *z-consistency* (Definitions 13 & 14), which might allow certain latent variables to remain entangled. Our results thus have the following flavor: Given a specific ground-truth causal graph G over z and a , we describe precisely the structure of the entanglement between latent factors via what we call an *entanglement graph* (Definition 3) and *graph preserving functions* (Definition 12). See Figure 3 for examples. Interestingly, the stronger identifiability up to permutation and element-wise transformation arises as a simple consequence of our theory when the graphical criterion of Lachapelle et al. (2022) is assumed to hold. In addition to these two main axes of generalization, we provide extensive examples illustrating the scope of our framework, our assumptions, and the consequences of our results (see Table 2 for a list). When it comes to the learning algorithm, we replace the sparsity *penalty* by a sparsity *constraint*, which improves the learning dynamics and is more interpretable, which results in easier hyperparameter tuning.

The hypothesis that *high-level concepts are related by a sparse dependency graph* has been described and leveraged for out-of-distribution generalization by Bengio (2019) and Goyal et al. (2021b), which were early sources of inspiration for this work. To the best of our knowledge, the theory developed in the present work and in its preliminary version (Lachapelle et al., 2022) is the

	Parametric assumption	Continuous a	Discrete a (interventions)	Temporal dependencies	Sufficient influence	Identifiable up to	Examples
Thm. 1	None	Required	–	Optional	Ass. 6	Def. 13	3, 4, 5, 9, 10
Thm. 2	None	–	Required	Optional	Ass. 7	Def. 13	3, 4, 5, 11, 12, 13
Thm. 3	None	Optional	Optional	Required	Ass. 8	Def. 14	6, 7, 8, 14
Thm. 4	Exp. fam.	Optional	Optional	Optional	Ass. 10	Def. 17	15
Thm. 5	Exp. fam.	Optional	Optional	Required	Ass. 11	Def. 14, 17	16

Table 1: Summary of our identifiability results.

first to show formally that this inductive bias can sometimes be enough to recover ground-truth latent factors.

Figure 1 shows a minimal motivating example in which our approach could be used to extract the high-level variables (such as the x -position of the three objects) and learn their dynamics (how the objects move and affect each other) from a time series of images and agent actions, $(\mathbf{x}^t, \mathbf{a}^t)$. Theorems 1, 2, 3 & 5 show how the sparse dependencies between the objects and the action can be leveraged to identify the latent variables and the graph describing their dynamics. The learned CGM could be used subsequently to simulate interventions on semantic variables (Pearl, 2009; Peters et al., 2017), such as changing the torque of the robot or the weight of the ball. Moreover, disentanglement could be useful to interpret what caused the actions of an agent (Pearl, 2019). Following Lachapelle et al. (2022), empirical works demonstrated that disentangled representations with sparse mechanisms can adapt to unseen interventions faster in the context of single-cell biology (Lopez et al., 2023) and synthetic video data (Lei et al., 2023).

Summary of our contributions:

1. We introduce¹ a new principle for disentanglement based on *mechanism sparsity regularization* motivated by rigorous and novel *identifiability guarantees* (Theorems 1, 2, 3 & 5).
2. We extend Lachapelle et al. (2022) by providing a *nonparametric* treatment and allowing for *arbitrary latent graphs*, as long as it has no edges between latent factors from the same time step (no instantaneous causal links). Given a latent ground-truth graph, our theory predicts the structure of the entanglement between variables, which we formalize with *entanglement graphs* (Definition 3), *graph preserving maps* (Definition 12) and novel *equivalence relations* (Definitions 13 & 14).
3. We provide several examples to illustrate the generality of our results and to get a better understanding of their various assumptions and consequences (summarized in Table 2). For instance, we show how multi-node interventions with unknown-targets can yield disentanglement, both with and without temporal dependencies (Examples 12 & 13).
4. We introduce an evaluation metric denoted by R_{con} that quantifies how close two representations are to being $\mathbf{a}$ -consistent or $\mathbf{z}$ -consistent (Section 6).
5. We implement a learning approach based on variational autoencoders (VAEs) (Kingma and Welling, 2014) that learns the mixing function $\mathbf{f}$, the transition distribution $p(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t})$ and the causal graph $\mathbf{G}$. The latter is learned using binary masks and regularized for sparsity using a *constraint* as opposed to a penalty as in Lachapelle et al. (2022).
6. We perform experiments on synthetic datasets to validate the prediction of our theory.

Overview. Section 2 introduces the model (Section 2.1), entanglement maps and graphs (Section 2.2), the notion of identifiability (Section 2.3), equivalence up to diffeomorphism (Section 2.4)

Examples	Type of disentanglement	Auxiliary variable	Time dependencies
3	Complete	Yes (single target)	Optional
4	Partial	Yes (single target)	Optional
5	Complete	Yes (multi-target)	Optional
6	Complete	Optional	Yes (independent factors)
7	Complete	Optional	Yes (dependent factors)
8	Partial	Optional	Yes (dependent factors)
9	Partial	Yes (single-target continuous)	Yes
10	Complete	Yes (multi-target continuous)	No
11	Complete	Yes (single-target interventions)	No
12	Complete	Yes (multi-target interventions)	Yes
13	Complete	Yes (grouped multi-target interventions)	No
14	Complete	No	Yes (non-Markovian)
16	Complete	No	Yes (Markovian)

Table 2: List of examples illustrating the scope of our theory, its assumptions and its consequences.

and disentanglement formally (Section 2.5). Section 3 provides mathematical intuition for why mechanism sparsity yields disentanglement (Section 3.2); introduces the machinery of graph preserving maps (Section 3.3) which are key to establish identifiability up to α -consistency (Section 3.4) and z -consistency (Section 3.5), i.e. *partial* disentanglement. Section 3 also discusses the relationship with interventions (Section 3.4.1), provides a graphical criterion guaranteeing *complete* disentanglement (Section 3.7), and introduces and discusses extensively the *sufficient influence assumptions* on which these results are critically based (Sections 3.8 & 3.9). Section 4 draws connections between our *nonparametric* theory and the *exponential family* assumption sometimes used in the literature. Section 5 presents the VAE-based learning algorithm with sparsity constraint. Section 6 introduces our novel R_{con} metric. Section 7 reviews the literature on identifiability in representation learning. Section 8 presents the empirical results.

Notation. Scalars are denoted in lower-case and vectors in lower-case bold, e.g. $x \in \mathbb{R}$ and $\mathbf{x} \in \mathbb{R}^n$. Note that these will sometimes denote random variables depending on the context. We maintain an analogous notation for scalar-valued and vector-valued functions, e.g., f and $\mathbf{f}$. The i th coordinate of the vector $\mathbf{x}$ is denoted by x_i . The set containing the first n integers excluding 0 is denoted by $[n]$. Given a subset of indices $S \subseteq [n]$, $\mathbf{x}_S$ denotes the subvector consisting of entries x_i for $i \in S$. Given a sequence of T random vectors $(\mathbf{x}^1, \dots, \mathbf{x}^T)$, the subsequence consisting of the first t elements is denoted by $\mathbf{x}^{\leq t} := (\mathbf{x}^1, \dots, \mathbf{x}^t)$, and analogously for $\mathbf{x}^{< t}$. We will sometimes combine these notation to get $\mathbf{x}_S^{\leq t} := (\mathbf{x}_S^1, \dots, \mathbf{x}_S^t)$. Given a function $\mathbf{f} : \mathbb{R}^n \rightarrow \mathbb{R}^m$, its Jacobian matrix evaluated at $\mathbf{x} \in \mathbb{R}^n$ is denoted by $D\mathbf{f}(\mathbf{x}) \in \mathbb{R}^{m \times n}$. See Table 5 in the Appendix for more.

2 Problem Setting, Entanglement Graphs & Disentanglement

In this section, we introduce the latent variable model under consideration (Section 2.1), entanglement graphs (Section 2.2), identifiability and observational equivalence (Section 2.3), equivalence up to diffeomorphism (Section 2.4) as well as permutation equivalence (Section 2.5).

2.1 An Identifiable Latent Causal Model

We now specify the setting under consideration. Assume that we observe the realization of a sequence of d_x -dimensional random vectors $\{\mathbf{x}^t\}_{t=1}^T$ and a sequence of d_a -dimensional auxiliary

vectors $\{\mathbf{a}^t\}_{t=0}^{T-1}$. The coordinates of $\mathbf{a}^t$ are discrete or continuous and can potentially represent, for example, an action taken by an agent or the index of an intervention or environment (see Section 3.4.1). Observations $\{\mathbf{x}^t\}$ are assumed to be explained by a sequence of continuous random vectors d_z of hidden dimensions $\{\mathbf{z}^t\}_{t=1}^T$ using a ground-truth decoder function $\mathbf{f}$.

Assumption 1 (Observation model) *For all $t \in [T]$, the observations $\mathbf{x}^t$ are given by*

$$\mathbf{x}^t = \mathbf{f}(\mathbf{z}^t) + \mathbf{n}^t,$$

where $\mathbf{n}^t \sim \mathcal{N}(0, \sigma^2 \mathbf{I})$ are mutually independent across time and independent of all $\mathbf{z}^t$ and $\mathbf{a}^t$ with $\sigma^2 \geq 0$. Moreover, $d_z \leq d_x$ and $\mathbf{f} : \mathbb{R}^{d_z} \rightarrow \mathbb{R}^{d_x}$ is a diffeomorphism onto its image². Lastly, assume that $\mathbf{f}(\mathbb{R}^{d_z})$ is closed in $\mathbb{R}^{d_x}$.

Importantly, we suppose that each factor $\mathbf{z}_i^t$ contains interpretable information about the observation, e.g., for high-dimensional images, the coordinates $\mathbf{z}_i^t$ might be the position of an object, its color, or its orientation in space. This idea that there exists a *ground-truth decoder* $\mathbf{f}$ that captures the relationship between the so-called “natural factors of variations” and the observations $\mathbf{x}$ is of capital importance, since it is the very basis for a mathematical definition of disentanglement (Definition 7). Appendix D.1 discusses the implications of the diffeomorphism assumption (see also Mansouri et al. (2022)). We denote $\mathbf{z}^{\leq t} := [\mathbf{z}^1 \dots \mathbf{z}^t] \in \mathbb{R}^{d_z \times t}$ and analogously for $\mathbf{z}^{< t}$ and other random vectors.

In a similar spirit to previous works on nonlinear ICA (Hyvärinen et al., 2019; Khemakhem et al., 2020a), we assume that the latent factors $\mathbf{z}_i^t$ are conditionally independent given the past.

Assumption 2 (Conditionally independent latent factors) *The latent factors $\mathbf{z}_i^t$ are conditionally mutually independent given $\mathbf{z}^{< t}$ and $\mathbf{a}^{< t}$:*

$$p(\mathbf{z}^t \mid \mathbf{z}^{< t}, \mathbf{a}^{< t}) = \prod_{i=1}^{d_z} p(\mathbf{z}_i^t \mid \mathbf{z}^{< t}, \mathbf{a}^{< t}), \quad (1)$$

where $p(\mathbf{z}^t \mid \mathbf{z}^{< t}, \mathbf{a}^{< t})$ is a density function w.r.t. the Lebesgue measure on $\mathbb{R}^{d_z}$. We assume that the support of $p(\mathbf{z}_i^t \mid \mathbf{z}^{< t}, \mathbf{a}^{< t})$ is $\mathbb{R}$ for all $\mathbf{z}^{< t}$ and $\mathbf{a}^{< t}$. The support of $p(\mathbf{z}^t \mid \mathbf{z}^{< t}, \mathbf{a}^{< t})$ is thus given by $\mathbb{R}^{d_z}$. At $t = 1$, we set $p(\mathbf{z}^1 \mid \mathbf{z}^{< 1}, \mathbf{a}^{< 1}) = p(\mathbf{z}^1 \mid \mathbf{a}^0) = \prod_{i=1}^{d_z} p(\mathbf{z}_i^1 \mid \mathbf{a}^0)$.

We will refer to the l.h.s. of (1) as the *transition model* and to the factors $p(\mathbf{z}_i^t \mid \mathbf{z}^{< t}, \mathbf{a}^{< t})$ as *mechanisms*. In a causal setting, this conditional independence assumption states that latent variables cannot exert influence on other latent variables instantaneously. We believe this to be a reasonable assumption when the time steps are sampled at a sufficiently high rate since, at a fundamental level, causal effects cannot travel faster than the speed of light. Of course in many practical scenarios, a low sampling rate can entail instantaneous dependencies which, if too strong, might

2. A *diffeomorphism* is a C^1 bijection with a C^1 inverse. Generally, given a map $\mathbf{h} : A \rightarrow \mathbb{R}^m$ where $A \subseteq \mathbb{R}^n$, saying $\mathbf{h}$ is C^k is typically only well defined if A is an open set of $\mathbb{R}^n$. Throughout, if $A \subseteq \mathbb{R}^n$ is arbitrary (not necessarily open), we say $\mathbf{h}$ is C^k if there exists a C^k map $\tilde{\mathbf{h}} : U \rightarrow \mathbb{R}^m$ defined on an open set U of $\mathbb{R}^n$ containing A such that $\mathbf{h} = \tilde{\mathbf{h}}$ on A . Note that it is then meaningful for $\mathbf{f}^{-1} : \mathbf{f}(\mathbb{R}^{d_z}) \rightarrow \mathbb{R}^{d_z}$ to be C^1 even when $\mathbf{f}(\mathbb{R}^{d_z})$ is not open in $\mathbb{R}^{d_x}$. Moreover, it can be shown that $\mathbf{f} : \mathbb{R}^{d_z} \rightarrow \mathbb{R}^{d_x}$ is a diffeomorphism onto its image if $\mathbf{f}$ is a homeomorphism onto its image, i.e. continuous in both directions, and has a full rank Jacobian everywhere on its domain (Munkres, 1991, Sec. 23 & Thm. 24.1).

result in model misspecification. In Section 7, we discuss and contrast with causal representation learning methods which allow for instantaneous influence.

Notice that we do not assume the dynamical system to be *Markovian*³, i.e. the distribution over future states can depend on the whole history of latents and auxiliary variables $(\mathbf{z}^{<t}, \mathbf{a}^{<t})$. In addition, this model can represent *non-homogeneous* processes by taking the auxiliary variable $\mathbf{a}$ to be a time index (Hyvärinen et al., 2019).

We are going to describe the dependency structure of the latent and auxiliary variables over time using a *probabilistic directed graphical model* composed of two bipartite graphs, $\mathbf{G}^z \in \{0, 1\}^{d_z \times d_z}$, which relates $\mathbf{z}^{<t}$ to $\mathbf{z}^t$, and $\mathbf{G}^a \in \{0, 1\}^{d_z \times d_a}$, which relates $\mathbf{a}^{<t}$ to $\mathbf{z}^t$. A directed edge points from $\mathbf{z}_j^{<t}$ to $\mathbf{z}_i^t$ if and only if $G_{i,j}^z = 1$. Analogously, a directed edge points from $\mathbf{a}_\ell^{<t}$ to $\mathbf{z}_i^t$, if and only if $G_{i,\ell}^a = 1$. Figure 1 shows an example of such graphs together with its adjacency matrix $\mathbf{G} := [\mathbf{G}^z, \mathbf{G}^a]$. The following assumption specifies the relationship between these graphs and the transition model.

Assumption 3 (Transition model p follows $\mathbf{G}$) For every mechanism $i \in [d_z]$,

$$p(\mathbf{z}_i^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t}) = p(\mathbf{z}_i^t \mid \mathbf{z}_{\mathbf{Pa}_i^z}^{<t}, \mathbf{a}_{\mathbf{Pa}_i^a}^{<t}), \quad (2)$$

where $\mathbf{Pa}_i^z \subseteq [d_z]$ and $\mathbf{Pa}_i^a \subseteq [d_a]$ are the sets of parents of $\mathbf{z}_i^t$ in $\mathbf{G}^z$ and $\mathbf{G}^a$, respectively.

The graph $\mathbf{G}$ thus encodes a set of conditional independence statements about the latent and auxiliary variables.⁴ We will say that *mechanisms are sparse* when the graphs $\mathbf{G}^a$ and $\mathbf{G}^z$ are sparse.

This model has three components that need to be learned: (i) the decoder function $\mathbf{f}$, (ii) the transition model over latent variables p , and (iii) the dependency graph $\mathbf{G}$. We gather all these components into $\boldsymbol{\theta} := (\mathbf{f}, p, \mathbf{G})$. Everything else in the model, i.e. d_z and σ^2 , is assumed to be known. We assume that σ^2 is known here mainly for simplicity, since, when it is not, it can be identified as shown by Lachapelle et al. (2022, Appendix A.4.1), as long as $d_x > d_z$.

Notice how we have not specified any model for the auxiliary variable $\mathbf{a}^t$. We do not intend to do so in this work, as we are solely interested in modeling the *conditional* distribution of $\mathbf{x}^{\leq T}$ and $\mathbf{z}^{\leq T}$ given $\mathbf{a}^{<T}$. We denote by $\mathcal{A} \subseteq \mathbb{R}^{d_a}$ the set of possible values for the auxiliary variable $\mathbf{a}^t$. We thus have that, for all values of $\mathbf{a}^{<T} \in \mathcal{A}^T$, our model induces a conditional distribution

$$p(\mathbf{x}^{\leq T} \mid \mathbf{a}^{<T}) = \int \prod_{t=1}^T p(\mathbf{x}^t \mid \mathbf{z}^t) p(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t}) d\mathbf{z}^{\leq T},$$

where $p(\mathbf{x}^t \mid \mathbf{z}^t) = \mathcal{N}(\mathbf{x}^t; \mathbf{f}(\mathbf{z}^t), \sigma^2 \mathbf{I})$. We note that if $\sigma^2 = 0$, the conditional distribution of $\mathbf{x}^t$ given $\mathbf{z}^t$ is a Dirac centered at $\mathbf{f}(\mathbf{z}^t)$ and thus has no density w.r.t. the Lebesgue measure. Even if, in that case, the above integral does not make sense, the conditional distribution of $\mathbf{x}^{\leq T}$ given $\mathbf{a}^{<T}$ is still well-defined and all the results of this work still hold since none of the proofs requires $\sigma^2 > 0$.

3. We use the term “Markovian” in the stochastic processes sense, as in Markov chains, where the future is independent of the past given the present. We will use this meaning throughout. This is to be contrasted with the Markovian property in structural models which states that exogenous noises are mutually independent (Bareinboim et al., 2022).

4. This notion of graph is not standard since, typically, directed edges point from scalar variables to scalar variables, whereas here a directed edge points from $\mathbf{z}_j^{<t}$ (a vector) to $\mathbf{z}_i^t$ (a scalar). This can be understood as having all directed edges $\mathbf{z}_j^\tau \rightarrow \mathbf{z}_i^t$ for all $\tau < t$. This does not imply that all edges are necessarily “active”. We made this choice to simplify the presentation and hypothesize that analogous results could be derived for standard graphs.

A motivating example. Figure 1 represents a minimal example in which our theory applies. The environment consists of three objects: a tree, a robot and a ball with x -positions z_T^t , z_R^t and z_B^t , respectively. Together, they form the vector z^t of high-level latent variables, i.e., $z^t = (z_T^t, z_R^t, z_B^t)$. A remote controls the direction in which the wheels of the robot turn. The vector a^t records these actions, which might be taken by a human or an artificial agent trained to accomplish some goal. The only observations are the actions a^t and the images x^t representing the scene which is given by $x^t = f(z^t) + n^t$. The dynamics of the environment is governed by the transition model p , which, e.g., could be given by a Gaussian model of the form $p(z_i^t | z^{<t}, a^{<t}) = \mathcal{N}(z_i^t; \mu_i(z^{t-1}, a^{t-1}), \sigma_z^2)$. Plausible connectivity graphs G^z and G^a are shown in Figure 1 showing how the latent factors are related and how the controller affects them. For every object, its position at the time step t depends on its position at $t - 1$. The position of the tree, z_T^t , is not affected by anything, since neither the robot nor the ball can change its position. The robot, z_R^t , changes its position based both on the action a^{t-1} and the position of the tree z_T^{t-1} (in case of collision). The position of the ball, z_B^t , is affected by both the robot, which can move around it by running into it, and the tree, which can bounce. The key observations here are that (i) the different objects interact *sparingly* with one another and (ii) the action a^t affects very few objects (in this case, only one). The theorems of Section 3 show how one can leverage this sparsity for disentanglement.

2.2 Entanglement Maps & Entanglement Graphs

In this section, we define *entanglement maps*, which describe the functional relationship between the learned and ground-truth representations, and *entanglement graphs*, which describe their entanglement structure.

Definition 1 (Entanglement maps) Let f and $\tilde{f}$ be two diffeomorphisms from $\mathbb{R}^{d_z}$ to their images such that $f(\mathbb{R}^{d_z}) = \tilde{f}(\mathbb{R}^{d_z})$. The *entanglement map* of the pair $(f, \tilde{f})$ is given by

$$v := f^{-1} \circ \tilde{f}. \quad (3)$$

This map will be crucial throughout this work, especially in defining disentanglement. Intuitively, the entanglement map for a pair of decoders $(f, \tilde{f})$ translates the representation of one model to that of the other. In general, the entanglement maps of $(f, \tilde{f})$ and $(\tilde{f}, f)$ are different.

We now define the *dependency graph* of some function h so that each edge indicates that some input i influences some output j :

Definition 2 (Functional dependency graph) Let h be a function from $\mathbb{R}^n$ to $\mathbb{R}^m$. The *dependency graph* of h is a bipartite directed graph from $[n]$ to $[m]$ with adjacency matrix $H \in \{0, 1\}^{m \times n}$ such that

$$H_{i,j} = 0 \iff \text{There is a function } \bar{h} \text{ such that, for all } a \in \mathbb{R}^n, h_i(a) = \bar{h}_i(a_{-j}),$$

where a_{-j} is a with its j th coordinate removed.

Example 1 (Dependency graph of a linear map) Let $h(z) := Wz$ where $W \in \mathbb{R}^{m \times n}$ and let H be the dependency graph of h . Then $H_{i,j} = 0 \iff W_{i,j} = 0$.

We will be particularly interested in the dependency graph of the entanglement map $v := f^{-1} \circ \tilde{f}$, denoted by V .

Definition 3 (Entanglement graphs) Let $\mathbf{f}$ and $\tilde{\mathbf{f}}$ be two diffeomorphisms from $\mathbb{R}^{d_z}$ to their images such that $\mathbf{f}(\mathbb{R}^{d_z}) = \tilde{\mathbf{f}}(\mathbb{R}^{d_z})$. The **entanglement graph** of the pair $(\mathbf{f}, \tilde{\mathbf{f}})$ is the dependency graph (Definition 2) of their entanglement map $\mathbf{v} := \mathbf{f}^{-1} \circ \tilde{\mathbf{f}}$, which we denote $\mathbf{V} \in \{0, 1\}^{d_z \times d_z}$.

We now relate the dependency graph of a function to the zeros of its Jacobian matrix. A proof can be found in Appendix A.2.

Proposition 1 (Linking dependency graph and Jacobian) Let $\mathbf{h}$ be a C^1 function, i.e. continuously differentiable, from $\mathbb{R}^n$ to $\mathbb{R}^m$ and let $\mathbf{H}$ be its dependency graph (Definition 2). Then,

$$\mathbf{H}_{i,j} = 0 \iff \text{For all } \mathbf{a} \in \mathbb{R}^n, D\mathbf{h}(\mathbf{a})_{i,j} = 0. \quad (4)$$

The equivalence (4) can be seen as an equivalent definition of dependency graph for differentiable functions.

2.3 Identifiability and Observational Equivalence

To analyze formally whether a specific algorithm is expected to yield a disentangled representation, we will rely on the notion of *identifiability*. Before we define what we mean by identifiability, we need the notion of *observationally equivalent* models. Two models are observationally equivalent if both models represent the same distribution over observations. The following formalizes this definition.

Definition 4 (Observational equivalence) We say that two models $\boldsymbol{\theta} := (\mathbf{f}, p, \mathbf{G})$ and $\tilde{\boldsymbol{\theta}} := (\tilde{\mathbf{f}}, \tilde{p}, \tilde{\mathbf{G}})$ satisfying Assumption 1 are **observationally equivalent**, denoted $\boldsymbol{\theta} \sim_{\text{obs}} \tilde{\boldsymbol{\theta}}$, if and only if, for all $\mathbf{a}^{<T} \in \mathcal{A}^T$ and all $\mathbf{x}^{\leq T} \in \mathbb{R}^{d_x \times T}$,

$$p(\mathbf{x}^{\leq T} \mid \mathbf{a}^{<T}) = \tilde{p}(\mathbf{x}^{\leq T} \mid \mathbf{a}^{<T}).$$

Formally, we say that a parameter $\boldsymbol{\theta}$ is **identifiable up to some equivalence relation** $\sim$, when

$$\boldsymbol{\theta} \sim_{\text{obs}} \tilde{\boldsymbol{\theta}} \implies \boldsymbol{\theta} \sim \tilde{\boldsymbol{\theta}}. \quad (5)$$

This work is mainly concerned with proving statements of the above form by making assumptions both on $\boldsymbol{\theta}$ and $\tilde{\boldsymbol{\theta}}$. The stronger the assumptions on $\boldsymbol{\theta}$ and $\tilde{\boldsymbol{\theta}}$ are, the stronger the equivalence relation $\sim$ will be. The following sections present two equivalence relations over models, namely $\sim_{\text{diff}}$ and $\sim_{\text{perm}}$. We note that the equivalence relation $\sim_{\text{perm}}$ will help us formalize disentanglement.

Practically speaking, observational equivalence between the learned model $\hat{\boldsymbol{\theta}}$ and the ground-truth model $\boldsymbol{\theta}$ can be achieved via maximum likelihood estimation in the infinite data regime. Thus, identifiability results of the form of (5) guaranty that if the learned model is perfectly fitted to the data (assumed infinite), its parameter $\hat{\boldsymbol{\theta}}$ is $\sim$ -equivalent to that of the ground-truth model, $\boldsymbol{\theta}$.

2.4 Equivalence up to Diffeomorphism

We start by defining *equivalence up to diffeomorphism*. This equivalence relation is important since we will show later on that it is actually the same as observational equivalence and will thus be our first step in all our identifiability results. In what follows, we overload the notation and write $\mathbf{v}(\mathbf{z}^{<t}) := [\mathbf{v}(\mathbf{z}^1), \dots, \mathbf{v}(\mathbf{z}^{t-1})]$, and similarly for other functions.

Definition 5 (Equivalence up to diffeomorphism) We say that two models $\theta := (\mathbf{f}, p, \mathbf{G})$ and $\tilde{\theta} := (\tilde{\mathbf{f}}, \tilde{p}, \tilde{\mathbf{G}})$ satisfying Assumption 1 are **equivalent up to diffeomorphism**, denoted $\theta \sim_{\text{diff}} \tilde{\theta}$, if and only if $\mathbf{f}(\mathbb{R}^{d_z}) = \tilde{\mathbf{f}}(\mathbb{R}^{d_z})$ and, for all $t \in [T]$, all $\mathbf{a}^{<t} \in \mathcal{A}^t$ and all $\mathbf{z}^{\leq t} \in \mathbb{R}^{d_z \times t}$,

$$\tilde{p}(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t}) = p(\mathbf{v}(\mathbf{z}^t) \mid \mathbf{v}(\mathbf{z}^{<t}), \mathbf{a}^{<t}) |\det D\mathbf{v}(\mathbf{z}^t)|, \quad (6)$$

where $\mathbf{v} := \mathbf{f}^{-1} \circ \tilde{\mathbf{f}}$ (entanglement map) is a diffeomorphism and $D\mathbf{v}$ denotes its Jacobian matrix.

The fact that the relation $\sim_{\text{diff}}$ is indeed an *equivalence* comes from the fact that the set of diffeomorphisms from a set to itself forms a group under composition.

To better understand the above definition, let $\mathbf{z}^t := \mathbf{g}(\mathbf{z}^{<t}, \mathbf{a}^{<t}; \epsilon^t)$ and $\tilde{\mathbf{z}}^t := \tilde{\mathbf{g}}(\tilde{\mathbf{z}}^{<t}, \mathbf{a}^{<t}; \tilde{\epsilon}^t)$ where ϵ^t and $\tilde{\epsilon}^t$ are noise variables and $\mathbf{g}$ and $\tilde{\mathbf{g}}$ are functions such that random variables $\mathbf{z}^t$ and $\tilde{\mathbf{z}}^t$ have conditional densities given by $p(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t})$ and $\tilde{p}(\tilde{\mathbf{z}}^t \mid \tilde{\mathbf{z}}^{<t}, \mathbf{a}^{<t})$, respectively. Using the change-of-variable formula for densities, one can rewrite (6) as

$$\tilde{\mathbf{g}}(\tilde{\mathbf{z}}^{<t}, \mathbf{a}^{<t}; \tilde{\epsilon}^t) \stackrel{d}{=} \mathbf{v}^{-1} \circ \mathbf{g}(\mathbf{v}(\tilde{\mathbf{z}}^{<t}), \mathbf{a}^{<t}; \epsilon^t), \quad (7)$$

where “ $\stackrel{d}{=}$ ” denotes equality in distribution. This equation has a nice interpretation: applying the latent transition model $\tilde{\theta}$ to go from $(\tilde{\mathbf{z}}^{<t}, \mathbf{a}^{<t})$ to $\tilde{\mathbf{z}}^t$ is the same as first applying $\mathbf{v}$, then applying the latent transition model θ and finally applying $\mathbf{v}^{-1}$. Equation (7) is reminiscent of Ahuja et al. (2022a), in which the mechanism $\tilde{\mathbf{g}}$ would be called an *imitator* of $\mathbf{g}$. Ahuja et al. (2022a) showed that $\sim_{\text{obs}}$ and $\sim_{\text{diff}}$ are actually one and the same. For completeness, we present an analogous argument here. We start by showing that $\theta \sim_{\text{diff}} \tilde{\theta}$ implies $\theta \sim_{\text{obs}} \tilde{\theta}$.

$$\begin{aligned} p(\mathbf{x}^{\leq T} \mid \mathbf{a}^{<T}) &= \int \prod_{t=1}^T [p(\mathbf{x}^t \mid \mathbf{z}^t) p(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t})] d\mathbf{z}^{\leq T} \\ &= \int \prod_{t=1}^T [p(\mathbf{x}^t \mid \mathbf{v}(\mathbf{z}^t)) p(\mathbf{v}(\mathbf{z}^t) \mid \mathbf{v}(\mathbf{z}^{<t}), \mathbf{a}^{<t})] |\det D\mathbf{v}(\mathbf{z}^{\leq T})| d\mathbf{z}^{\leq T} \\ &= \int \prod_{t=1}^T [p(\mathbf{x}^t \mid \mathbf{v}(\mathbf{z}^t)) p(\mathbf{v}(\mathbf{z}^t) \mid \mathbf{v}(\mathbf{z}^{<t}), \mathbf{a}^{<t}) |\det D\mathbf{v}(\mathbf{z}^t)|] d\mathbf{z}^{\leq T} \\ &= \int \prod_{t=1}^T [\tilde{p}(\mathbf{x}^t \mid \mathbf{z}^t) \tilde{p}(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t})] d\mathbf{z}^{\leq T} = \tilde{p}(\mathbf{x}^{\leq T} \mid \mathbf{a}^{<T}), \end{aligned}$$

where the second equality used the change-of-variable formula, the third equality used the fact that the Jacobian of $\mathbf{v}(\mathbf{z}^{\leq T})$ is block-diagonal (each block corresponds to a time step t) and the next to last equality used the definition of $\sim_{\text{diff}}$ and the fact that

$$p(\mathbf{x}^t \mid \mathbf{v}(\mathbf{z}^t)) = \mathcal{N}(\mathbf{x}^t; \mathbf{f}(\mathbf{f}^{-1} \circ \tilde{\mathbf{f}}(\mathbf{z}^t)), \sigma^2 \mathbf{I}) = \mathcal{N}(\mathbf{x}^t; \tilde{\mathbf{f}}(\mathbf{z}^t), \sigma^2 \mathbf{I}) = \tilde{p}(\mathbf{x}^t \mid \mathbf{z}^t).$$

The following proposition establishes the converse, i.e. that $\theta \sim_{\text{obs}} \tilde{\theta}$ implies $\theta \sim_{\text{diff}} \tilde{\theta}$. Since its proof is more involved, we present it in the Appendix A.3. Note that this first identifiability result is relatively weak and should be seen as a first step towards stronger guarantees. A very similar result was shown by Ahuja et al. (2022a, Theorem 3.1) to highlight the fact that the representation $\mathbf{f}$ is identifiable up to the equivariances $\mathbf{v}$ of the transition model p .

Proposition 2 (Identifiability up to diffeomorphism) *Let $\theta := (\mathbf{f}, p, \mathbf{G})$ and $\hat{\theta} := (\hat{\mathbf{f}}, \hat{p}, \hat{\mathbf{G}})$ be two models satisfying Assumption 1. If $\theta \sim_{\text{obs}} \hat{\theta}$ (Def. 4), then $\theta \sim_{\text{diff}} \hat{\theta}$ (Def. 5).*

Intuitively, Proposition 2 shows that if two models agree on the distribution of the observations, then their “data manifold” $\mathbf{f}(\mathbb{R}^{d_z})$ and $\hat{\mathbf{f}}(\mathbb{R}^{d_z})$ are equal and their respective transition models are related via $\mathbf{v} := \mathbf{f}^{-1} \circ \hat{\mathbf{f}}$.

2.5 Disentanglement and Equivalence up to Permutation

A disentangled representation is often defined intuitively as a representation in which the coordinates are in one-to-one correspondence with *natural factors of variation* in the data. We are going to assume that these natural factors are captured by an unknown ground-truth decoder $\tilde{\mathbf{f}}$. Given a learned decoder $\hat{\mathbf{f}}$ such that $\mathbf{f}(\mathbb{R}^{d_z}) = \hat{\mathbf{f}}(\mathbb{R}^{d_z})$, the entanglement map $\mathbf{v} := \mathbf{f}^{-1} \circ \hat{\mathbf{f}}$ gives a correspondence between the learned representation $\hat{\mathbf{f}}$ and the natural factors of variations of $\mathbf{f}$. The following equivalence relation will help us define disentanglement.

Definition 6 (Equivalence up to permutation) *We say that two models $\theta := (\mathbf{f}, p, \mathbf{G})$ and $\tilde{\theta} := (\tilde{\mathbf{f}}, \tilde{p}, \tilde{\mathbf{G}})$ satisfying Assumptions 1, 2 & 3 are **equivalent up to permutation**, denoted $\theta \sim_{\text{perm}} \tilde{\theta}$, if and only if there exists a permutation matrix $\mathbf{P}$ such that*

1. $\theta \sim_{\text{diff}} \tilde{\theta}$ (Def. 5) and $\tilde{\mathbf{G}}^a = \mathbf{P}\mathbf{G}^a$ and $\tilde{\mathbf{G}}^z = \mathbf{P}\mathbf{G}^z\mathbf{P}^\top$; and
2. The entanglement map $\mathbf{v} := \mathbf{f}^{-1} \circ \tilde{\mathbf{f}}$ can be written as $\mathbf{v} = \mathbf{d} \circ \mathbf{P}^\top$, where $\mathbf{d}$ is element-wise, i.e. $d_i(\mathbf{z})$ depends only on z_i , for all i . In other words, the entanglement graph is $\mathbf{V} = \mathbf{P}^\top$.

The fact that the relation $\sim_{\text{perm}}$ is an equivalence relation is actually a special case of a more general result that we present later in Section 3.4.

This allows us to give a formal definition of (complete) disentanglement. Note that we use the term *complete* to contrast with *partial* disentanglement.

Definition 7 (Complete disentanglement) *Given a ground-truth model θ , we say that a learned model $\hat{\theta}$ is **completely disentangled** when $\theta \sim_{\text{perm}} \hat{\theta}$.*

Intuitively, a learned representation is *completely disentangled* when there is a one-to-one correspondence between its coordinates and those of the ground-truth representation (see Figure 2).

We define *partial* disentanglement as something that lives strictly between equivalence up to diffeomorphism and equivalence up to permutation:

Definition 8 (Partial disentanglement) *Given a ground-truth model θ , we say that a learned model $\hat{\theta}$ is **partially disentangled** when $\theta \sim_{\text{diff}} \hat{\theta}$ with an entanglement graph $\mathbf{V}$ (Definition 3) that is neither a permutation nor the complete graph.*

This definition of partial disentanglement ranges from models that are almost completely entangled, i.e. those with very dense entanglement graphs $\mathbf{V}$, to ones that are very close to being completely disentangled, i.e. those with a very sparse $\mathbf{V}$. The following section will make more precise how one can learn a completely or partially disentangled representation from data and exactly what form the entanglement graph is going to take.

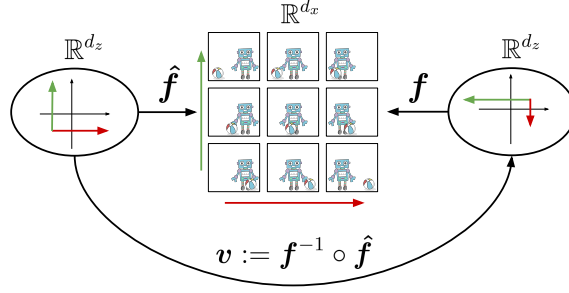


Figure 2: An illustration of disentanglement (Definition 7). The ground-truth decoder f captures the “natural factors of variations”, which here are the x -positions of the robot and ball. The learned decoder $\hat{f}$ is disentangled here because each of its latent coordinates corresponds exactly to one object in the scene, in a one-to-one fashion. Mathematically, this is captured by the special structure of the entanglement map $v := f^{-1} \circ \hat{f}$, which is a permutation composed with an element-wise invertible transformation.

3 Nonparametric Partial Disentanglement via Mechanism Sparsity

In this section, we show that without further assumptions the model introduced so far is unidentifiable (Section 3.1), provide a first insight as to why mechanism sparsity can make the model identifiable and thus lead to disentanglement (Section 3.2), introduce the machinery of G -preserving maps (Section 3.3) which leads up to theorems showing identifiability up to α -consistency (Section 3.4) and z -consistency (Section 3.5), which corresponds to partial disentanglement. We also relate these results to interventions (Section 3.4.1), show how to combine both regularization on $\hat{G}^a$ and $\hat{G}^z$ to obtain stronger guarantees (Section 3.6) and introduce a graphical criterion guaranteeing complete disentanglement (Section 3.7). Finally, we introduce the *sufficient influence assumptions* and prove the identifiability results (Section 3.8), and provide multiple examples to build intuition (Section 3.9).

3.1 Conditional Independence is Insufficient for Disentanglement and Graph Discovery

In the setting where the sequence $\{z^t\}_{t=1}^T$ is fully observed, the absence of instantaneous causal effects (Assumption 2) combined with faithfulness (when the only conditional independencies in the distribution are those implied by the causal graph) makes the causal graph fully identifiable (Peters et al., 2017, Section 10.3.1). In our setting where only the sequence $\{f(z^t) + n^t\}_{t=1}^T$ is observed, conditional independence alone is insufficient to identify the latent causal graph or the latent variables, as the following example demonstrates.

Example 2 Suppose $p(z^t \mid z^{<t}, a^{<t}) := \mathcal{N}(z^t; Wz^{t-1} + a^{t-1}e_1, \eta^2 I)$ where $a^t \in \mathcal{A} = \mathbb{R}$, $W \in \mathbb{R}^{d \times d}$ is symmetric and $e_1 \in \mathbb{R}^d$ is the first one-hot vector. Let G^z and G^a be the graph entailed by this model. Note that G^z can be very dense while G^a contains only a single edge. Choose f to be any diffeomorphism on its image. This defines a model $\theta = (f, p, G)$ that satisfies the conditional independence assumption. We now construct a second model $\tilde{\theta}$ that is observationally equivalent to θ , but has a drastically different graph $\tilde{G}$ and different latent factors. Since W is symmetric, it can be diagonalized by an orthogonal matrix, i.e., $W = UDU^\top$, where U is orthogonal and D

is diagonal. Define $\tilde{\mathbf{f}} := \mathbf{f} \circ \mathbf{U}$, so that $\mathbf{f}(\mathbb{R}^d) = \tilde{\mathbf{f}}(\mathbb{R}^d)$ and $\mathbf{v} := \mathbf{f}^{-1} \circ \tilde{\mathbf{f}} = \mathbf{U}$. Define

$$\tilde{p}(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t}) := p(\mathbf{U}\mathbf{z}^t \mid \mathbf{U}\mathbf{z}^{<t}, \mathbf{a}^{<t}),$$

which is a well-defined conditional density by the change-of-variable formula (using $|\det \mathbf{U}| = 1$). We can see $\tilde{p}$ as the conditional density resulting from first applying $\mathbf{U}$ to $\mathbf{z}^{t-1}$, taking a step using p to get $\mathbf{z}^t$ and finally transforming $\mathbf{z}^t$ using $\mathbf{U}^\top$, analogously to (7). It is clear that $\tilde{p}$ satisfies the conditional independence (Assumption 2) since an orthogonal transformation of an isotropic Gaussian remains an isotropic Gaussian. Furthermore, the mean of $\tilde{p}(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t})$ is given by $\mathbf{U}^\top \mathbf{W} \mathbf{U} \mathbf{z}^{t-1} + \mathbf{a}^{t-1} \mathbf{U}^\top \mathbf{e}_1 = \mathbf{D} \mathbf{z}^{t-1} + \mathbf{a}^{t-1} (\mathbf{U}^\top)_{:,1}$. We find that $\tilde{p}$ is Markov w.r.t. $\tilde{\mathbf{G}}^z = \mathbf{I}$, since $\mathbf{D}$ is diagonal, and $\tilde{\mathbf{G}}^a := \mathbb{1}$, assuming $\mathbf{U}$ is dense. Note that $\tilde{\mathbf{G}}^z$ will be in general much sparser than $\mathbf{G}^z$ and $\tilde{\mathbf{G}}^a$ is potentially much denser than $\mathbf{G}^a$ (how extreme these differences are depends on the exact choice of $\mathbf{W}$). Moreover, in general, the representation is not disentangled, since $\mathbf{v} = \mathbf{U}$ can be chosen to be dense. Nevertheless, these two models are observationally equivalent since they are equivalent up to diffeomorphism by construction (Section 2.4).

The preceding example illustrates some difficulties one might face when it comes to the identifiability of the model introduced so far. This serves as a motivation to explore additional inductive biases such as mechanism sparsity, which we do next.

3.2 A First Mathematical Insight for Disentanglement via Mechanism Sparsity

In this section, we derive a first insight pointing towards how mechanism sparsity regularization, i.e., regularizing $\hat{\mathbf{G}}$ to be sparse, can promote disentanglement.

Before going further, we briefly introduce an abuse of notation that will be handy throughout: we will sometimes use vectors and matrices as sets of indices corresponding to their supports.

Definition 9 (Vectors & matrices as index sets) Let $\mathbf{a} \in \mathbb{R}^n$ and $\mathbf{A} \in \mathbb{R}^{m \times n}$. We will sometimes use $\mathbf{a}$ to denote the set of indices corresponding to the support of the vector $\mathbf{a}$, i.e.

$$\mathbf{a} \sim \{i \in [n] \mid \mathbf{a}_i \neq 0\}.$$

This will allow us to write things like $i \in \mathbf{a}$ or $\mathbf{a} \subseteq \mathbf{b}$, where $\mathbf{b} \in \mathbb{R}^n$. We will use an analogous convention for matrices, i.e.,

$$\mathbf{A} \sim \{(i, j) \in [m] \times [n] \mid \mathbf{A}_{i,j} \neq 0\},$$

This will allow us to write things like $(i, j) \in \mathbf{A}$ and $\mathbf{A} \subseteq \mathbf{B}$, where $\mathbf{B} \in \mathbb{R}^{m \times n}$.

Recall that we would like to show that $\boldsymbol{\theta} \sim_{\text{obs}} \hat{\boldsymbol{\theta}}$ implies $\boldsymbol{\theta} \sim_{\text{perm}} \hat{\boldsymbol{\theta}}$, i.e., disentanglement (or partial disentanglement). Our approach will be to start from (6), which is guaranteed by Proposition 2, and perform a series of algebraic manipulations to gain mathematical insight into how regularizing $\hat{\mathbf{G}}$ to be sparse (mechanism sparsity) can induce disentanglement. A key manipulation will be taking first and second order derivatives. For this to be possible, we require a certain level of smoothness for the transition models:

Assumption 4 (Smoothness of transition model) When $\mathbf{a}$ is continuous, the transition densities $p(\mathbf{z}_i^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t})$ are C^2 functions from $\mathbb{R} \times \mathbb{R}^{d_z \times (t-1)} \times \mathcal{A}^t$ to $\mathbb{R}$ and $\mathcal{A} \subseteq \mathbb{R}^\ell$ is regular closed⁵. When $\mathbf{a}$ is discrete (e.g. Section 3.4.1), for all $\mathbf{a}^{<t}$, $p(\mathbf{z}_i^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t})$ are C^2 functions from $\mathbb{R} \times \mathbb{R}^{d_z \times (t-1)}$ to $\mathbb{R}$.

5. A set $\mathcal{A} \subseteq \mathbb{R}^\ell$ is regular closed when it is equal to the closure of its interior, i.e. $\overline{\mathcal{A}^\circ} = \mathcal{A}$.

We start by taking the log on both sides of (6) and let $q := \log p$ and $\hat{q} := \log \hat{p}$:

$$\hat{q}(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t}) = q(\mathbf{v}(\mathbf{z}^t) \mid \mathbf{v}(\mathbf{z}^{<t}), \mathbf{a}^{<t}) + \log |\det D\mathbf{v}(\mathbf{z}^t)|.$$

We then take the derivative w.r.t. $\mathbf{z}^t$ on both sides:

$$D_{\mathbf{z}}^t \hat{q}(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t}) = D_{\mathbf{z}}^t q(\mathbf{v}(\mathbf{z}^t) \mid \mathbf{v}(\mathbf{z}^{<t}), \mathbf{a}^{<t}) D\mathbf{v}(\mathbf{z}^t) + \eta(\mathbf{z}^t) \in \mathbb{R}^{1 \times d_z}, \quad (8)$$

where $D_{\mathbf{z}}^t q$ denotes the Jacobian of $q(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t})$ w.r.t. $\mathbf{z}^t$ and analogously for $D_{\mathbf{z}}^t \hat{q}$. The term $\eta(\mathbf{z}^t)$ is the derivative of $\log |\det D\mathbf{v}(\mathbf{z}^t)|$ w.r.t. $\mathbf{z}^t$.

We differentiate⁶ yet once more w.r.t. $\mathbf{a}^\tau$ for some $\tau < t$ (assuming $\mathbf{a}^t$ is continuous for now) and obtain

$$H_{z,a}^{t,\tau} \hat{q}(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t}) = D\mathbf{v}(\mathbf{z}^t)^\top H_{z,a}^{t,\tau} q(\mathbf{v}(\mathbf{z}^t) \mid \mathbf{v}(\mathbf{z}^{<t}), \mathbf{a}^{<t}) \in \mathbb{R}^{d_z \times d_a}, \quad (9)$$

where $H_{z,a}^{t,\tau} q \in \mathbb{R}^{d_z \times d_a}$ is the Hessian matrix of second derivatives w.r.t. $\mathbf{z}^t$ and $\mathbf{a}^\tau$ and similarly for $H_{z,a}^{t,\tau} \hat{q}$.

We now look more closely at some specific entry (i, ℓ) of the Hessian $H_{z,a}^{t,\tau} q$. We first see that

$$\frac{\partial^2}{\partial \mathbf{a}_\ell^\tau \partial \mathbf{z}_i^t} q(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t}) = \frac{\partial^2}{\partial \mathbf{a}_\ell^\tau \partial \mathbf{z}_i^t} \sum_{j=1}^{d_z} q(\mathbf{z}_j^t \mid \mathbf{z}_{\mathbf{Pa}_j^z}^{<t}, \mathbf{a}_{\mathbf{Pa}_j^a}^{<t}) \quad (10)$$

$$= \frac{\partial}{\partial \mathbf{a}_\ell^\tau} \sum_{j=1}^{d_z} \frac{\partial}{\partial \mathbf{z}_i^t} q(\mathbf{z}_j^t \mid \mathbf{z}_{\mathbf{Pa}_j^z}^{<t}, \mathbf{a}_{\mathbf{Pa}_j^a}^{<t}) \quad (11)$$

$$= \frac{\partial}{\partial \mathbf{a}_\ell^\tau} \frac{\partial}{\partial \mathbf{z}_i^t} q(\mathbf{z}_i^t \mid \mathbf{z}_{\mathbf{Pa}_i^z}^{<t}, \mathbf{a}_{\mathbf{Pa}_i^a}^{<t}), \quad (12)$$

where the first equality holds by (1) & (2) and a basic property of logarithms. It is clear that (12) equals zero when $\ell \notin \mathbf{Pa}_i^a$. This is a crucial observation, since it implies that whenever $\mathbf{G}_{i,\ell}^a = 0$, we also have $(H_{z,a}^{t,\tau} q)_{i,\ell} = 0$. In other words, $H_{z,a}^{t,\tau} q \subseteq \mathbf{G}^a$. The same argument can also be applied to get $H_{z,a}^{t,\tau} \hat{q} \subseteq \hat{\mathbf{G}}^a$. Notice how this argument would fail without the conditional independence assumption, since $\mathbf{z}_i^t$ could appear in terms other than $q(\mathbf{z}_i^t \mid \mathbf{z}_{\mathbf{Pa}_i^z}^{<t}, \mathbf{a}_{\mathbf{Pa}_i^a}^{<t})$, as it could be the parent of some other $\mathbf{z}_j^t$, for $j \neq i$.

Intuitive argument. We can start to see why regularizing $\hat{\mathbf{G}}$ to be sparse might induce disentanglement. Intuitively, a sparse $\hat{\mathbf{G}}^a$ forces $D\mathbf{v}(\mathbf{z}^t)$ to be sparse, since otherwise the l.h.s. of (13) will not be sparse:

$$\underbrace{H_{z,a}^{t,\tau} \hat{q}(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t})}_{\subseteq \hat{\mathbf{G}}^a} = \underbrace{D\mathbf{v}(\mathbf{z}^t)^\top}_{\text{forced to be sparse}} \underbrace{H_{z,a}^{t,\tau} q(\mathbf{v}(\mathbf{z}^t) \mid \mathbf{v}(\mathbf{z}^{<t}), \mathbf{a}^{<t})}_{\subseteq \mathbf{G}^a}, \quad (13)$$

And, of course, the sparser $D\mathbf{v}(\mathbf{z}^t)$ is, the more disentangled $\hat{\mathbf{f}}$ is, since $D\mathbf{v}_{i,j} = 0$ everywhere implies $\mathbf{V}_{i,j} = 0$ under weak assumptions (Proposition 1). The above argument is not rigorous and is provided only to build intuition. It will be formalized later.

6. This derivative is well defined on $\mathcal{A}$ (in the sense that it does not depend on its C^k extension) since $\mathcal{A}$ is regular close. We prove this general fact in Lemma 4 in the appendix.

Sparse temporal dependencies. In what precedes, we made use of the sparsity of the graph $\hat{G}^a$ to argue that Dv must also be sparse. We now show a similar intuition based on the sparsity of $\hat{G}^z$. Starting from (8), instead of differentiating w.r.t. $\mathbf{a}^\tau$, we will differentiate w.r.t. $\mathbf{z}^\tau$, for some $\tau < t$, which yields:

$$H_{z,z}^{t,\tau} \hat{q}(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t}) = Dv(\mathbf{z}^t)^\top H_{z,z}^{t,\tau} q(\mathbf{v}(\mathbf{z}^t) \mid \mathbf{v}(\mathbf{z}^{<t}), \mathbf{a}^{<t}) Dv(\mathbf{z}^\tau) \in \mathbb{R}^{d_z \times d_z}, \quad (14)$$

where $H_{z,z}^{t,\tau} q$ is the Hessian matrix of second derivatives of q w.r.t. $\mathbf{z}^t$ and $\mathbf{z}^\tau$, and analogously for $H_{z,z}^{t,\tau} \hat{q}$. Using an argument perfectly analogous to Equations (10) to (12), we can show that whenever $G_{i,j}^z = 0$, we also have $(H_{z,z}^{t,\tau} q)_{i,j} = 0$, and similarly for $\hat{G}^z$ and $H_{z,z}^{t,\tau} \hat{q}$. In other words, $H_{z,z}^{t,\tau} q \subseteq G^z$ and $H_{z,z}^{t,\tau} \hat{q} \subseteq \hat{G}^z$. Therefore, analogously to (13), regularizing $\hat{G}^z$ to be sparse intuitively should force Dv to be sparse as well, i.e., bringing us closer to disentanglement:

$$\underbrace{H_{z,z}^{t,\tau} \hat{q}(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t})}_{\subseteq \hat{G}^z} = \underbrace{Dv(\mathbf{z}^t)^\top}_{\text{forced to be sparse}} \underbrace{H_{z,z}^{t,\tau} q(\mathbf{v}(\mathbf{z}^t) \mid \mathbf{v}(\mathbf{z}^{<t}), \mathbf{a}^{<t})}_{\subseteq G^z} \underbrace{Dv(\mathbf{z}^\tau)}_{\text{forced to be sparse}}. \quad (15)$$

The crux of our technical contribution in this work is to make the above arguments formal and precisely characterize what the sparsity structure of $Dv(\mathbf{z})$ (hence, $\mathbf{V}$) will be based on the ground-truth graph G (Theorems 1, 2 & 3). We also provide conditions on G to guarantee complete disentanglement (Proposition 7).

3.3 Graph Preserving Maps

Theorems 1, 2, 3 & 5 will show how regularizing $\hat{G}$ to be sparse can force the dependency graph of the entanglement map v to be sparse as well. These results characterize the functional dependency structure of the entanglement map v as a function of the ground-truth graph G . This link will be made precise thanks to the notion of graph preserving maps, which we define next. Before going further, we need to set up the following notation.

Definition 10 (Aligned subspaces of $\mathbb{R}^m$ and $\mathbb{R}^{m \times n}$) Given a binary vector $\mathbf{b} \in \{0, 1\}^m$, let

$$\mathbb{R}_{\mathbf{b}}^m := \{\mathbf{x} \in \mathbb{R}^m \mid \mathbf{b}_i = 0 \implies \mathbf{x}_i = 0\}$$

Given a binary matrix $\mathbf{B} \in \{0, 1\}^{m \times n}$, let

$$\mathbb{R}_{\mathbf{B}}^{m \times n} := \{\mathbf{M} \in \mathbb{R}^{m \times n} \mid \mathbf{B}_{i,j} = 0 \implies \mathbf{M}_{i,j} = 0\}.$$

Note that $\mathbb{R}_{\mathbf{b}}^m$ and $\mathbb{R}_{\mathbf{B}}^{m \times n}$ are vector spaces under addition. This means that given $\mathbf{a}^{(1)}, \dots, \mathbf{a}^{(k)} \in \mathbb{R}_{\mathbf{b}}^m$, we have that $\text{span}\{\mathbf{a}^{(1)}, \dots, \mathbf{a}^{(k)}\} \subseteq \mathbb{R}_{\mathbf{b}}^m$, where span denotes the subspace of all linear combinations. Similarly, given $\mathbf{A}^{(1)}, \dots, \mathbf{A}^{(k)} \in \mathbb{R}_{\mathbf{B}}^{m \times n}$, we have that $\text{span}\{\mathbf{A}^{(1)}, \dots, \mathbf{A}^{(k)}\} \subseteq \mathbb{R}_{\mathbf{B}}^{m \times n}$.

To start formally reasoning about what will happen if $\hat{G}^a$ is regularized sparse, we temporarily assume that $\hat{G}^a = G^a$. With this assumption, we can interpret (13) as implying that $Dv(\mathbf{z}^t)^\top$ must *preserve* the “sparsity structure” of the matrix $H_{z,a}^{t,\tau} q$. This observation motivates the following definitions, which will be central to our contribution.

Definition 11 (G -preserving matrix) Given $G \in \{0, 1\}^{m \times n}$, a matrix $\mathbf{C} \in \mathbb{R}^{m \times m}$ is G -preserving when

$$\mathbf{C}^\top \mathbb{R}_G^{m \times n} \subseteq \mathbb{R}_G^{m \times n}.$$

Definition 12 (G -preserving functions) Given $G \in \{0, 1\}^{m \times n}$, a function $c : \mathbb{R}^m \rightarrow \mathbb{R}^m$ is G -preserving when its dependency graph C (Definition 2) is G -preserving.

Without surprise, a linear map $c(z) := Cz$ where $C \in \mathbb{R}^{m \times m}$ is G -preserving (Definition 12) if and only if the matrix C is G -preserving (Definition 11).

We now show that G -preserving functions can be defined alternatively in terms of a simple condition on their dependency graph. This characterization of G -preserving functions is key to understanding how (partial) disentanglement results from sparsity regularization.

Proposition 3 A function c with dependency graph C (Definition 2) is G -preserving if and only

$$G_{i,\cdot} \not\subseteq G_{j,\cdot} \implies C_{i,j} = 0, \text{ for all } i, j.$$

Proof We start by showing the “only if” statement. We suppose $G_{i,\cdot} \not\subseteq G_{j,\cdot}$ and must now show that $C_{i,j} = 0$. We know that there exists k such that $G_{i,k} = 1$ but $G_{j,k} = 0$. Since $C^\top \mathbb{R}_G^{m \times n} \subseteq \mathbb{R}_G^{m \times n}$ and $e_i e_k^\top \in \mathbb{R}_G^{m \times n}$, we must have $C^\top e_i e_k^\top \in \mathbb{R}_G^{m \times n}$. Since $G_{j,k} = 0$, we must have $0 = (C^\top e_i e_k^\top)_{j,k} = C_{i,j}$.

We now show the “if” statement. Let $A \in \mathbb{R}_G^{m \times n}$. Take some (i, j) such that $G_{i,j} = 0$. We must now show that $(C^\top A)_{i,j} = 0$. We have $(C^\top A)_{i,j} = \sum_k C_{k,i} A_{k,j}$. We now check that each term in this sum must be zero. If $A_{k,j} = 0$, of course the corresponding term is zero. If $A_{k,j} \neq 0$, it implies that $G_{k,j} = 1$ and thus $G_{k,\cdot} \not\subseteq G_{i,\cdot}$. By assumption, this implies that $C_{k,i} = 0$ and thus $C_{k,i} A_{k,j} = 0$. Hence $(C^\top A)_{i,j} = 0$ as desired. ■

We now characterize differentiable G -preserving functions in terms of their Jacobian matrices.

Lemma 1 A differentiable function $c : \mathbb{R}^m \rightarrow \mathbb{R}^m$ is G -preserving if and only if, for all $z \in \mathbb{R}^m$, $Dc(z)$ is G -preserving.

Proof Assume c is G -preserving with dependency graph C . By Proposition 3, this is equivalent to having that, for all $i, j \in [n]$,

$$G_{i,\cdot} \not\subseteq G_{j,\cdot} \implies C_{i,j} = 0.$$

But by Proposition 1, this statement is equivalent to

$$G_{i,\cdot} \not\subseteq G_{j,\cdot} \implies \forall z \in \mathbb{R}^m, Dc(z)_{i,j} = 0,$$

which is equivalent to saying that $Dc(z)$ is G -preserving for all $z \in \mathbb{R}^m$ (again by Proposition 3). ■

We will now show that G -preserving diffeomorphisms form a group under composition. To do so, we start by showing that invertible G -preserving matrices form a group under matrix multiplication (Proposition 4) and extend the result to diffeomorphisms in Proposition 5.

Proposition 4 Invertible G -preserving matrices form a group under matrix multiplication.

Proof We must show that the set of invertible G -preserving matrices contains the identity, is closed under matrix multiplication, and is closed under inversion.

Clearly, I is G -preserving since $I^\top \mathbb{R}_G^{m \times n} = \mathbb{R}_G^{m \times n}$.

Let C_1 and C_2 be G -preserving. Then, $C_1 C_2$ is G -preserving because

$$(C_1 C_2)^\top \mathbb{R}_G^{m \times n} = C_2^\top C_1^\top \mathbb{R}_G^{m \times n} \subseteq C_2^\top \mathbb{R}_G^{m \times n} \subseteq \mathbb{R}_G^{m \times n}.$$

Let C be G -preserving and invertible. Since $C^\top$ is invertible as a map from $\mathbb{R}^{m \times n}$ to $\mathbb{R}^{m \times n}$, the dimensionality of the subspace $\mathbb{R}_G^{m \times n}$ must be equal to the dimensionality of $C^\top \mathbb{R}_G^{m \times n}$. This fact combined with $C^\top \mathbb{R}_G^{m \times n} \subseteq \mathbb{R}_G^{m \times n}$ implies that $C^\top \mathbb{R}_G^{m \times n} = \mathbb{R}_G^{m \times n}$. Hence $\mathbb{R}_G^{m \times n} = (C^{-1})^\top \mathbb{R}_G^{m \times n}$, i.e. C^{-1} is G -preserving. $\blacksquare$

We now extend the above results to diffeomorphisms using Proposition 1.

Proposition 5 *The set of G -preserving diffeomorphisms forms a group under composition.*

Proof We must show that the set of G -preserving diffeomorphisms contains the identity, is closed under matrix multiplication, and is closed under inversion.

The first statement is trivial since the entanglement graph of the identity diffeomorphism is the identity graph $C := I$, and of course it is G -preserving.

We now prove the second statement. Let c and c' be two diffeomorphisms with dependency graph C and C' respectively. By the chain rule, we have that

$$D(c \circ c')(z) = Dc(c'(z))Dc'(z).$$

By Lemma 1, we have that $Dc(c'(z))$ and $Dc'(z)$ are G -preserving matrices and, by Proposition 4 their product must also be G -preserving. Hence $D(c \circ c')(z)$ is G -preserving for all z and thus, by Lemma 1, $c \circ c'$ is G -preserving.

The proof of the third statement has a similar flavor. By the inverse function theorem, we have

$$Dc^{-1}(z) = Dc(c^{-1}(z))^{-1}.$$

Moreover, by Lemma 1, $Dc(c^{-1}(z))$ is G preserve. Furthermore, its inverse is also G -preserving by Proposition 4. Similarly to the previous step, because c^{-1} is C^1 , we can use Lemma 1 to conclude that c^{-1} is also G -preserving. $\blacksquare$

3.4 Nonparametric Identifiability via Auxiliary Variables with Sparse Influence

In this section, we introduce our first identifiability results based on the sparsity of the graph G^a which describes the structure of the dependencies between $\mathbf{a}^{<t}$ and $\mathbf{z}^t$. We will see that, under some assumptions, regularizing the learned graph $\hat{G}^a$ to be sparse will allow identifiability up to the following equivalence class:

Definition 13 (α -consistency equivalence) *We say two models $\theta := (\mathbf{f}, p, G)$ and $\tilde{\theta} := (\tilde{\mathbf{f}}, \tilde{p}, \tilde{G})$ satisfying Assumptions 1, 2 & 3 are α -consistent, denoted $\theta \sim_{\text{con}}^\alpha \tilde{\theta}$, if and only if there exists a permutation matrix P such that*

1. $\theta \sim_{\text{diff}} \tilde{\theta}$ (Def. 5), and $\tilde{G}^a = PG^a$; and
2. the entanglement map $v := f^{-1} \circ \tilde{f}$ can be written as $v = c \circ P^\top$ where c is a G^a -preserving diffeomorphism (Def. 12).

The main difference between α -consistency (above definition) and permutation equivalence (Definition 6), is that instead of having $v = d \circ P^\top$ where d is element-wise, we have $v = c \circ P^\top$ where c is G^a -preserving, which allows some mixing between the latent factors. Importantly, a G^a -preserving map typically has missing edges in its dependency graph, as Proposition 3 shows. This means that this equivalence relation imposes structure on the entanglement map v . Depending on the structure of G^a , this can mean either complete, partial, or no disentanglement whatsoever. Note that the equivalence $\sim_{\text{perm}}$ is stronger than $\sim_{\text{con}}^a$, in the sense that $\theta \sim_{\text{perm}} \hat{\theta} \implies \theta \sim_{\text{con}}^a \hat{\theta}$. This is because element-wise transformations d are always G -preserving, for any G .

We demonstrate in Appendix A.4 that the α -consistency relation is indeed an equivalence relation, as claimed in the above definition. This follows from the fact that the set of G^a -preserving diffeomorphisms forms a *group* under composition (Proposition 5).

The first result provides conditions under which regularizing the learned graph $\hat{G}^a$ to be as sparse as the ground-truth graph G^a will induce the learned model to be α -consistent with the ground-truth one.

Theorem 1 (Nonparametric disentanglement from continuous α with sparse influence) *Let the parameters $\theta := (f, p, G)$ and $\hat{\theta} := (\hat{f}, \hat{p}, \hat{G})$ correspond to two models satisfying Assumptions 1, 2, 3, & 4. Further assume that*

1. **[Observational equivalence]** $\theta \sim_{\text{obs}} \hat{\theta}$ (Def. 4);
2. **[Sufficient influence of α]** The Hessian matrix $H_{z,a}^{t,\tau} \log p(z^t \mid z^{<t}, \alpha^{<t})$ varies “sufficiently”, as formalized in Assumption 6;

Then, there exists a permutation matrix P such that $PG^a \subseteq \hat{G}^a$. Further assume that

3. **[Sparsity regularization]** $\|\hat{G}^a\|_0 \leq \|G^a\|_0$;

Then, $\theta \sim_{\text{con}}^a \hat{\theta}$ (Def. 13).

The second assumption as well as a proof of this result is delayed to Section 3.8 for pedagogical reasons. We now describe and provide intuition about each assumption one by one.

Observational equivalence. The first assumption simply requires that both models agree on the observational model. In practice, this is achieved by fitting the model to data.

Sufficient influence. The second assumption requires that the “effect” of $\alpha^{<t}$ on z^t is “sufficiently strong”. The assumption will be formalized and discussed in more detail later in Sections 3.8 & 3.9, but we can already see that it concerns the Hessian matrix $H_{z,a}^{t,\tau} \log p$ that we saw earlier in Eq. (13) of Sect. 3.2.

Sparsity regularization. The first two assumptions imply that the learned graph $\hat{G}^a$ is a super-graph of some permutation of the ground-truth graph G^a . By adding the *sparsity regularization* assumption, we find that the learned graph $\hat{G}^a$ is *exactly* a permutation of the ground-truth graph G^a and that, more precisely, the learned model is $\sim_{\text{con}}^a$ -equivalent to the ground-truth. This assumption is satisfied if $\hat{G}^a$ is a minimal graph among all graphs that allow the model to exactly match the ground-truth generative distribution. In Section 5, we suggest achieving this in practice by adding a sparsity penalty in the training objective or by constraining the optimization problem.

α -consistency. The final conclusion of the result states that the learned model is $\sim_{\text{con}}^a$ -equivalent to the ground-truth, which means the entanglement map $v := f^{-1} \circ \hat{f}$ can be written as $v = c \circ P^\top$ where c is G^a -preserving. This is important since the G^a -preserving condition imposes structure on the entanglement graph V (Definition 3), as implied by Proposition 3. In other words, the result predicts precisely which latent factors are expected to remain entangled.

Remark 1 (Inverse of v) We defined v to be the mapping from the learned to the ground-truth representation, but in some context, it might be more telling to look at v^{-1} , which maps from the ground-truth to the learned representation. If $v = c \circ P^\top$ where c is G^a -preserving (as predicted by Theorem 1), we know that its inverse is given by $v^{-1} = P \circ c^{-1}$ where c^{-1} is G^a -preserving, by closure under inversion (Proposition 5).

The following result is the same as above, but for *discrete* auxiliary variables a . This case is very important to cover the case where a indexes sparse interventions targeting latent factors, which we discuss in more detail in Section 3.4.1. Note that the only difference from the above theorem is the “sufficient influence” assumption, which we will present formally in Section 3.8 together with a proof of the result.

Theorem 2 (Nonparametric disentanglement via discrete a with sparse influence) *Let the parameters $\theta := (f, p, G)$ and $\hat{\theta} := (\hat{f}, \hat{p}, \hat{G})$ correspond to two models satisfying Assumptions 1, 2, 3 & 4. Further assume that*

1. **[Observational equivalence]** $\theta \sim_{\text{obs}} \hat{\theta}$ (Def. 4);
2. **[Sufficient influence of a]** *The vector of derivatives $D_z^t \log p(z^t \mid z^{<t}, a^{<t})$ depends “sufficiently strongly” on each component a_ℓ , as formalized in Assumption 7;*

Then, there exists a permutation matrix P such that $PG^a \subseteq \hat{G}^a$. Further assume that

3. **[Sparsity regularization]** $\|\hat{G}^a\|_0 \leq \|G^a\|_0$;

Then, $\theta \sim_{\text{con}}^a \hat{\theta}$ (Def. 13).

We now provide a few examples to illustrate how Theorems 1 & 2 can be applied. Here, we concentrate on the relationship between the graph G^a and the entanglement graph V (Definition 3). The question of whether or not the sufficient influence assumption is satisfied will be delayed to Section 3.9, where the examples will be made more concrete by specifying latent models more explicitly.

Example 3 ($G^a = I$ implies complete disentanglement) Assume that $d_a = d_z$ and $G^a = I$, i.e., each latent variable is affected by only one auxiliary variable, and each auxiliary variable affects only one latent variable. The graph G^a is depicted in Figure 3a and G^z could be anything (see the remark below). Assuming that the ground-truth transition model satisfies the sufficient influence assumption of Theorem 1 or 2, we have $\theta \sim_{\text{obs}} \hat{\theta}$ & $\|\hat{G}^a\|_0 \leq \|G^a\|_0 \implies \theta \sim_{\text{con}}^a \hat{\theta}$. This means that there exists a permutation matrix P such that $\hat{G}^a = PG^a$ and such that the entanglement map is given by $v = c \circ P^\top$ where c is a G^a -preserving diffeomorphism (Definition 11). But since $G^a = I$, Proposition 3 tells us that the dependency graph of c is simply $C := I$ and thus the entanglement graph is $V = P^\top$, i.e., complete disentanglement holds. In fact, one could add more columns to G^a (i.e. adding auxiliary variables) without changing the conclusion. Example 11 will provide a concrete example satisfying the sufficient influence assumption of Theorem 2.

Remark 2 (Temporal dependencies are not necessary) The above example did not mention anything about the temporal graph G^z . That is because this graph could be anything, in fact, we could be in the special case where there are no temporal dependencies whatsoever, i.e., $T = 1$ and the latent model is simply $p(z \mid a) = \prod_{i=1}^{d_z} p(z_i \mid a)$. In that case, Theorems 1 & 2 could still be applied to prove the identifiability of the representation, as long as their assumptions hold. This remark also applies to the next two examples.

Example 4 (Action targeting a single latent variable identifies it) Consider the situation depicted in Figure 1 where z_1 is the tree position, z_2 is the robot position and z_3 is the ball position ($d_z = 3$). Assume $a \in \mathbb{R}$ corresponds to the torque applied to the robot wheels ($d_a = 1$). We therefore have $G^a = [0, 1, 0]^\top$, i.e., a affects only z_2 . For the sake of this example, G^z can be anything, i.e., it does not have to be lower triangular like in Figure 1 (see remark above).

If the sufficient influence assumption of Theorem 1 or 2 is satisfied, we see that $\theta \sim_{\text{obs}} \hat{\theta}$ & $\|\hat{G}^a\|_0 \leq \|G^a\|_0$ implies $v = c \circ P^\top$ where P is a permutation and c is a G^a -preserving diffeomorphism. Using Proposition 3, this means that the dependency graph of c is given by

$$C = \begin{bmatrix} * & * & * \\ 0 & * & 0 \\ * & * & * \end{bmatrix}, \text{ since } G_{2,\cdot}^a \not\subseteq G_{1,\cdot}^a \text{ and } G_{2,\cdot}^a \not\subseteq G_{3,\cdot}^a,$$

where “*” indicates a potentially nonzero value. This means that one of the components of the learned representation will be an invertible transformation of the ground-truth variable z_2 (robot position), while the other components could be a mixture of z_1 , z_2 and z_3 . Figure 3b shows both the graph G^a and the corresponding entanglement graph V assuming $P = I$. Example 9 will make this example more concrete by explicitly specifying a latent model that satisfies the sufficient influence assumption of Theorem 1.

Example 5 (Complete disentanglement from multi-target actions) Assume $d_z = 3$ and $d_a = 3$ where $G^a \in \mathbb{R}^{d_z \times d_a}$ is given by Figure 3c and the temporal graph G^z could be anything (see Remark 2 above). If the sufficient influence assumption of Theorem 1 or 2 is satisfied, then we have that $\theta \sim_{\text{obs}} \hat{\theta}$ & $\|\hat{G}^a\|_0 \leq \|G^a\|_0$ implies $v = c \circ P^\top$ where P is a permutation and c is a G^a -preserving diffeomorphism. Proposition 3 implies that the dependency graph of c is simply $C := I$ because $G_{i,\cdot}^a \not\subseteq G_{j,\cdot}^a$ for all distinct i, j . This means that we have complete disentanglement (Definition 7). Examples 10, 12 and 13 will explore more concrete instantiations of this example by specifying concrete latent models satisfying the sufficient influence assumptions of Theorems 1 and 2.

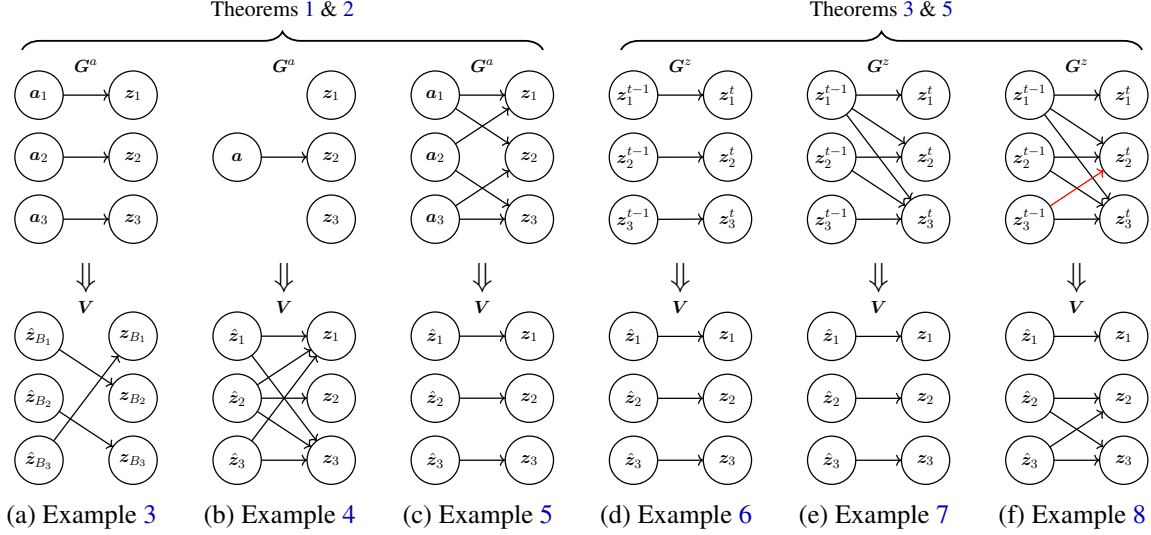


Figure 3: Graphs G^a and G^z from Examples 3, 4, 5, 6, 7 & 8 with their respective entanglement graphs V (Definition 3) guaranteed by Theorems 1, 2, 3 & 5 (assuming $P = I$ for simplicity). Recall, that V describes the dependency structure of $v = f^{-1} \circ \hat{f}$, which maps $\hat{z}$ to z . By Remark 1, the functional dependency graph of v^{-1} is exactly the same except for z and $\hat{z}$ being interchanged.

3.4.1 UNKNOWN-TARGET INTERVENTIONS ON THE LATENT FACTORS

An important special case of Theorem 2 is when $\mathbf{a}^{t-1}$ corresponds to a one-hot vector indexing an *intervention with unknown targets* on the latent variables $\mathbf{z}^t$. This specific kind of intervention has been explored previously in the context of causal discovery where the intervention occurs on *observed* variables instead of *latent* variables (Eaton and Murphy, 2007; Mooij et al., 2020; Squires et al., 2020; Jaber et al., 2020; Brouillard et al., 2020; Ke et al., 2019). Recently, multiple works in causal representation learning have considered interventions on latent variables (Lachapelle et al., 2022; Lippe et al., 2023b; Ahuja et al., 2023; Squires et al., 2023; Buchholz et al., 2023; von Kügelgen et al., 2023; Zhang et al., 2023; Jiang and Aragam, 2023) (see Section 7 for more). Here is how our framework can accommodate such interventions: Assume $\mathbf{a}^{t-1} \in \{\vec{0}, \mathbf{e}_1, \dots, \mathbf{e}_{d_a}\}$, where each $\mathbf{e}_\ell$ is a one-hot vector. The action $\mathbf{a}^{t-1} = \vec{0}$ corresponds to the *observational setting*, i.e., when no intervention occurred, while $\mathbf{a}^{t-1} = \mathbf{e}_\ell$ corresponds to the ℓ th intervention. In that context, the unknown graph G^a describes which latents are targeted by the intervention, i.e. $\ell \in \mathbf{Pa}_i^a$ if and only if z_i is targeted by the ℓ th intervention. To see this, recall that, under Assumption 3, we have

$$p(z_i^t | \mathbf{z}^{<t}, \mathbf{a}^{<t}) = p(z_i^t | \mathbf{z}_{\mathbf{Pa}_i^z}^{<t}, \mathbf{a}_{\mathbf{Pa}_i^a}^{t-1}),$$

where we implicitly assumed that $p(z_i^t | \mathbf{z}^{<t}, \mathbf{a}^{<t})$ does not depend on $\mathbf{a}^{<t-1}$. In the observational setting, i.e., when $\mathbf{a}^{t-1} = \vec{0}$, the conditional becomes $p(z_i^t | \mathbf{z}^{<t}, \vec{0})$. Now suppose that we are in the ℓ th intervention, i.e., $\mathbf{a}^{t-1} = \mathbf{e}_\ell$. Then, if $\ell \notin \mathbf{Pa}_i^a$, we have that $\mathbf{a}_{\mathbf{Pa}_i^a}^{t-1} = \vec{0}$, which means that the conditional is also $p(z_i^t | \mathbf{z}^{<t}, \vec{0})$, meaning that the variable z_i^t is *not* targeted by the ℓ th intervention. When $\ell \in \mathbf{Pa}_i^a$, we have $\mathbf{a}_{\mathbf{Pa}_i^a}^{t-1} \neq \vec{0}$ and thus the conditional is allowed to change freely, i.e., z_i^t is targeted by the ℓ th intervention.

Importantly, the assumption that $\mathbf{G}^a$ is sparse corresponds precisely to the *sparse mechanism shift* hypothesis from Schölkopf et al. (2021), i.e. that *only a few mechanisms change at a time*. Thm. 2 thus provides precise conditions for when sparse mechanism shifts induce disentanglement. Interestingly, our theory covers both hard and soft interventions, as long as the sufficient influence assumption is satisfied.

Remark 3 (Examples revisited) *Examples 3, 4 and 5 can be revisited while keeping in mind the “unknown-target intervention interpretation” in which $\mathbf{G}^a$ describes which latent variable is targeted by each intervention. For example, Example 3 tells us that if each latent variable is targeted by a single-node intervention, then complete disentanglement is guaranteed. Examples 11, 12 and 13 provide mathematically concrete latent models where $\mathbf{a}$ is interpreted to be an intervention.*

Remark 4 (Causal representation learning without temporal dependencies (static)) *The special case where $T = 1$, i.e., no temporal dependencies, is of special interest. In that case, the latent variable model is simply $p(\mathbf{z} \mid \mathbf{a}) = \prod_{i=1}^{d_z} p(z_i \mid \mathbf{a})$. In other words, the causal graph that relates the latent factors z_i is empty. In contrast, recent work on learning causal representation showed how to obtain disentanglement in general latent causal graphical models without temporal dependencies, but are limited to single-node interventions (Ahuja et al., 2023; Squires et al., 2023; Buchholz et al., 2023; von Kügelgen et al., 2023; Zhang et al., 2023; Jiang and Aragam, 2023). These works can be roughly thought of as fitting our interventional setting described in this section where (i) $T = 1$; (ii) conditional independence (Assumption 1) is dropped (allowing instantaneous causal effects); and (iii) only single-node interventions are allowed, corresponding to a very sparse $\mathbf{G}^a$ with columns containing at most one non-zero element, like in Example 3. Although our framework with $T = 1$ assumes that the causal graph between latent variables is empty, it allows for multi-node interventions which are sometimes sufficient to disentangle (Example 13). See Section 3.9.2 for more on this.*

3.5 Nonparametric Identifiability via Sparse Temporal Dependencies

This section is analogous to the previous one, but instead of leveraging the sparsity of $\mathbf{G}^a$ to show identifiability, it leverages the sparsity of $\mathbf{G}^z$, which describes the structure of the dependencies between the latents from one time step to another. We will see that, under some assumptions, regularizing the learned graph $\hat{\mathbf{G}}^z$ to be sparse will allow identifiability up to the following equivalence class:

Definition 14 (z-consistency equivalence) *We say two models $\theta := (\mathbf{f}, p, \mathbf{G})$ and $\tilde{\theta} := (\tilde{\mathbf{f}}, \tilde{p}, \tilde{\mathbf{G}})$ satisfying Assumptions 1, 2 & 3 are **z-consistent**, denoted $\theta \sim_{\text{con}}^z \tilde{\theta}$, if and only if there exists a permutation matrix $\mathbf{P}$ such that*

1. $\theta \sim_{\text{diff}} \tilde{\theta}$ (Def. 5) and $\tilde{\mathbf{G}}^z = \mathbf{P} \mathbf{G}^z \mathbf{P}^\top$; and
2. the entanglement map $\mathbf{v} := \mathbf{f}^{-1} \circ \tilde{\mathbf{f}}$ can be written as $\mathbf{v} = \mathbf{c} \circ \mathbf{P}^\top$ where $\mathbf{c}$ is a $\mathbf{G}^z$ -preserving and $(\mathbf{G}^z)^\top$ -preserving diffeomorphism (Definition 12).

This relation can be shown to be an *equivalence* relation, as was the case for $\sim_{\text{con}}^a$. This is shown in Appendix A.4. Analogously to $\sim_{\text{con}}^a$, the equivalence relation $\sim_{\text{con}}^z$ relates the structure of the entanglement map $\mathbf{v}$ to the graph $\mathbf{G}^z$ via the notion of $\mathbf{G}$ -preserving maps. It is also true that $\theta \sim_{\text{perm}} \hat{\theta} \implies \theta \sim_{\text{con}}^z \hat{\theta}$.

The following result is analogous to Theorems 1 and 2 where, instead of regularizing $\hat{G}^a$ to be sparse, we regularize $\hat{G}^z$. The next theorem shows how this type of sparsity regularization can induce the learned model to be z -consistent with the ground-truth one.

Theorem 3 (Nonparametric disentanglement via sparse temporal dependencies) *Let the parameters $\theta := (\mathbf{f}, p, \mathbf{G})$ and $\hat{\theta} := (\hat{\mathbf{f}}, \hat{p}, \hat{\mathbf{G}})$ correspond to two models satisfying Assumptions 1, 2, 3 & 4. Further assume that*

1. *[Observational equivalence] $\theta \sim_{\text{obs}} \hat{\theta}$ (Def. 4);*
2. *[Sufficient influence of z] The Hessian matrix $H_{z,z}^{t,\tau} \log p(z^t \mid z^{<t}, \alpha^{<t})$ varies “sufficiently”, as formalized in Assumption 8;*

Then, there exists a permutation matrix $\mathbf{P}$ such that $\mathbf{P}\mathbf{G}^z\mathbf{P}^\top \subseteq \hat{\mathbf{G}}^z$. Further assume that

3. *[Sparsity regularization] $\|\hat{\mathbf{G}}^z\|_0 \leq \|\mathbf{G}^z\|_0$;*

Then, $\theta \sim_{\text{con}}^z \hat{\theta}$ (Def. 14).

The structure of the above theorem is very similar to Theorem 1 & 2. For example, we still have a “sufficient influence” condition, but this time it concerns the Hessian matrix $H_{z,z}^{t,\tau} \log p$ which we saw in Section 3.2, Equation (15). The conclusion is that both models will be z -consistent, which means that we recover the graph $\mathbf{G}^z$ up to permutation and have that the entanglement map $\mathbf{v}$ has a dependency graph given by $\mathbf{V} = \mathbf{C}\mathbf{P}^\top$ where $\mathbf{C}$ is $\mathbf{G}^z$ - and $(\mathbf{G}^z)^\top$ -preserving. Section 3.8 introduces the sufficient influence assumption formally as well as a proof of Theorem 3.

We now build intuition through some minimal examples that show how one can apply the above theorem to draw links between the graph $\mathbf{G}^z$ and the resulting entanglement graph $\mathbf{V}$ (Definition 3). For now we simply assume that the assumption of sufficient influence (Assumption 8) is satisfied and wait until Section 3.9.3 to present more concrete transition models that satisfy it.

Example 6 (Disentanglement via independent factors with temporal dependencies) *Consider the situation depicted in Figure 3d where the graph $\mathbf{G}^z = \mathbf{I}$, i.e., the latents $\mathbf{z}_i^t$ are dependent in time but independent across dimensions. For this example, actions are unnecessary. Assuming the sufficient influence assumption of Theorem 3 is satisfied, we have $\theta \sim_{\text{obs}} \hat{\theta}$ & $\|\hat{\mathbf{G}}^z\|_0 \leq \|\mathbf{G}^z\|_0 \implies \theta \sim_{\text{con}}^z \hat{\theta}$, which means there exists a permutation $\mathbf{P}$ such that $\hat{\mathbf{G}}^z = \mathbf{P}\mathbf{G}^z\mathbf{P}^\top$ and such that the entanglement map is given by $\mathbf{v} = \mathbf{c} \circ \mathbf{P}^\top$ where $\mathbf{c}$ is $\mathbf{G}^z$ - and $(\mathbf{G}^z)^\top$ -preserving. Using Proposition 3, one can verify that the dependency graph of $\mathbf{c}$ is $\mathbf{C} = \mathbf{I}$ and thus $\mathbf{V} = \mathbf{P}^\top$, i.e. the learned representation is completely disentangled. Example 14 will provide a concrete transition model where the sufficient influence assumption of Theorem 3 holds for this simple graph $\mathbf{G}^z$.*

Example 7 (Disentanglement via sparsely dependent factors with temporal dependencies) *The previous examples assumed independent latents, i.e., $\mathbf{G}^z = \mathbf{I}$. Instead, we now consider a more interesting “lower triangular” graph $\mathbf{G}^z$, as depicted in Figures 3e (This is the same graph as in the tree-robot-ball example in Figure 1). Again using Proposition 3, one can verify that $\mathbf{C} = \mathbf{I}$ and thus $\mathbf{V} = \mathbf{P}^\top$, i.e. the learned representation is completely disentangled. Example 14 will provide a concrete transition model where the sufficient influence assumption of Theorem 3 holds.*

Example 8 (Partial disentanglement via temporal sparsity) Assume the same situation as previously, but add an additional edge from z_B^{t-1} to z_R^t (see Figure 3f). This could occur, for example, if the robot tries to follow the ball and is thus influenced by it. Using Proposition 3, one can show that c being G^z - and $(G^z)^\top$ -preserving means that its dependency graph is given by

$$C = \begin{bmatrix} * & 0 & 0 \\ 0 & * & * \\ 0 & * & * \end{bmatrix}.$$

This means the robot and the ball remain entangled in the learned representation.

3.6 Combining Sparsity Regularization on $\hat{G}^a$ & $\hat{G}^z$

A natural question at this point is whether Theorem 1 (or Theorem 2) can be combined with Theorem 3 to obtain stronger guarantees. The answer is yes. In this section, we explain how this can be done. We would like to show how the combination of the assumptions of Theorem 1 and Theorem 3 can yield identifiability up to the following stronger equivalence relation.

Definition 15 ((a, z)-consistency equivalence) We say two models $\theta := (f, p, G)$ and $\tilde{\theta} := (\tilde{f}, \tilde{p}, \tilde{G})$ satisfying Assumptions 1, 2 & 3 are (a, z)-**consistent**, denoted $\theta \sim_{\text{con}}^{z,a} \tilde{\theta}$, if and only if there exists a permutation matrix P such that

1. $\theta \sim_{\text{diff}} \tilde{\theta}$ (Def. 5) and $\tilde{G}^a = P^\top G^a$ and $\tilde{G}^z = P^\top G^z P$; and
2. the entanglement map $v := f^{-1} \circ \tilde{f}$ can be written as $v = c \circ P^\top$ where c is a G^a -, G^z - and $(G^z)^\top$ -preserving diffeomorphism (Def. 12).

Of course, if the assumptions of both theorems hold, we must have $\theta \sim_{\text{con}}^a \hat{\theta}$ and $\theta \sim_{\text{con}}^z \hat{\theta}$. As one might guess, this implies $\theta \sim_{\text{con}}^{a,z} \hat{\theta}$, as the following proposition shows. The reason why this result is not completely trivial is that the permutations P given by $\sim_{\text{con}}^a$ and $\sim_{\text{con}}^z$ might not be the same. Its proof can be found in Appendix A.4.1.

Proposition 6 Let $\theta := (f, p, G)$ and $\tilde{\theta} := (\tilde{f}, \tilde{p}, \tilde{G})$ be two models satisfying Assumptions 1, 2 & 3. We have $\theta \sim_{\text{con}}^{z,a} \tilde{\theta}$ if and only if $\theta \sim_{\text{con}}^a \tilde{\theta}$ and $\theta \sim_{\text{con}}^z \tilde{\theta}$.

We can thus combine both Theorems 1 (or Theorem 2) with Theorem 3 to obtain stronger guarantees. Practically, this means that regularizing both $\hat{G}^a$ and $\hat{G}^z$ to be sparse will lead to a more disentangled representation, i.e. a sparser entanglement graph V , than if regularization was applied only on $\hat{G}^a$ or only on $\hat{G}^z$.

3.7 Graphical Criterion for Complete Disentanglement

The previous sections introduced results guaranteeing identifiability up to $\sim_{\text{con}}^a$, $\sim_{\text{con}}^z$ and $\sim_{\text{con}}^{z,a}$, which all correspond to potentially *partial* disentanglement. This section provides an additional assumption to guarantee identifiability up to $\sim_{\text{perm}}$, i.e., *complete* disentanglement.

One can easily see from the definitions that $\theta \sim_{\text{perm}} \tilde{\theta}$ holds precisely when $\theta \sim_{\text{con}}^{a,z} \tilde{\theta}$ with $C = I$. This condition can be achieved by making an additional assumption on G . This assumption is taken directly from Lachapelle et al. (2022).

Assumption 5 (Graphical criterion, Lachapelle et al. (2022)) Let $G = [G^z \ G^a]$ be a graph. For all $i \in \{1, \dots, d_z\}$,

$$\left(\bigcap_{j \in \text{Ch}_i^z} \text{Pa}_j^z \right) \cap \left(\bigcap_{j \in \text{Pa}_i^z} \text{Ch}_j^z \right) \cap \left(\bigcap_{\ell \in \text{Pa}_i^a} \text{Ch}_\ell^a \right) = \{i\},$$

where Pa_i^z and Ch_i^z are the sets of parents and children of node z_i in G^z , respectively, while Ch_ℓ^a is the set of children of a_ℓ in G^a .

The following proposition shows that when G satisfies the above criterion, the set of models that are $\sim_{\text{con}}^{a,z}$ -equivalent to θ is equal to the set of models that are $\sim_{\text{perm}}$ -equivalent to θ , thus allowing complete disentanglement. See Appendix A.6 for a proof.

Proposition 7 (Complete disentanglement as a special case) Let $\theta := (f, p, G)$ and $\hat{\theta} := (\hat{f}, \hat{p}, \hat{G})$ be two models satisfying Assumptions 1, 2 & 3. If $\theta \sim_{\text{con}}^{z,a} \hat{\theta}$ and G satisfies Assumption 5, then $\theta \sim_{\text{perm}} \hat{\theta}$.

The above result shows that our general theory can guarantee complete disentanglement as a special case. This is one way in which our work generalizes the work of Lachapelle et al. (2022), in addition to relaxing the exponential family assumption. The following section explores how the exponential family assumption fits into our nonparametric theory and how it allows one to simplify the “sufficient influence assumptions”. But before, we provide some example to illustrate when Assumption 5 holds.

For example, the graphical criterion of Assumption 5 is trivially satisfied when G^z is diagonal, since $\{i\} = \text{Pa}_i^z$ for all i (actions are not necessary here). This simple case amounts to having mutual independence between the sequences $z_i^{\leq T}$, which is a standard assumption in the ICA literature (Tong et al., 1990; Hyvarinen and Morioka, 2017; Klindt et al., 2021). The illustrative example we introduced in Fig. 1 has a more interesting “non-diagonal” graph satisfying our criterion. Indeed, we have $\{T\} = \text{Pa}_T^z$, $\{R\} = \text{Ch}_R^z \cap \text{Pa}_R^z$, and $\{B\} = \text{Ch}_B^z$. This example is actually part of an interesting family of graphs that satisfy our criterion:

Proposition 8 (Sufficient condition for the graphical criterion) If $G_{i,i}^z = 1$ for all i (all nodes have a self-loop) and G^z has no 2-cycles, then G satisfies Assumption 5.

Proof Self-loops guarantee $i \in \text{Pa}_i^z \cap \text{Ch}_i^z$ for all i . Suppose $j \in \text{Pa}_i^z \cap \text{Ch}_i^z$ for some $i \neq j$. This implies that i and j form a 2-cycle, which is a contradiction. Thus, $\{i\} = \text{Pa}_i^z \cap \text{Ch}_i^z$ for all i . ■

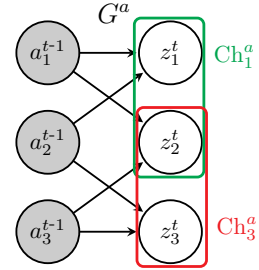


Figure 4: An example satisfying Assumption 5. Indeed, $\{z_1\} = \text{Ch}_1^a \cap \text{Ch}_2^a$, $\{z_2\} = \text{Ch}_1^a \cap \text{Ch}_3^a$ and $\{z_3\} = \text{Ch}_2^a \cap \text{Ch}_3^a$.

3.8 Proofs of Theorems 1, 2 & 3 and their Sufficient Influence Assumptions

In this section, we introduce the sufficient influence assumptions and use them to prove Theorems 1, 2 & 3. In the next section (Section 3.9), we provide multiple examples to gain intuition about the sufficient influence assumptions. Throughout, the following lemma will come in handy.

Lemma 2 (Invertible matrix contains a permutation) *Let $\mathbf{L} \in \mathbb{R}^{m \times m}$ be an invertible matrix. Then, there exists a permutation σ such that $\mathbf{L}_{i,\sigma(i)} \neq 0$ for all i , or, in other words, $\mathbf{P}^\top \subseteq \mathbf{L}$ where $\mathbf{P}$ is the permutation matrix associated with σ , i.e., $\mathbf{P}\mathbf{e}_i = \mathbf{e}_{\sigma(i)}$. Note that this implies that $\mathbf{P}\mathbf{L}$ and $\mathbf{L}\mathbf{P}$ have no zero on their diagonals.*

Proof Since the matrix $\mathbf{L}$ is invertible, its determinant is non-zero, i.e.

$$\det(\mathbf{L}) := \sum_{\sigma \in \mathfrak{S}_m} \text{sign}(\sigma) \prod_{i=1}^m \mathbf{L}_{i,\sigma(i)} \neq 0,$$

where $\mathfrak{S}_m$ is the set of m -permutations. This equation implies that at least one term of the sum is non-zero, meaning there exists a permutation σ such that, for all i , $\mathbf{L}_{i,\sigma(i)} \neq 0$. $\blacksquare$

3.8.1 SUFFICIENT INFLUENCE ASSUMPTION OF THEOREM 1 AND ITS PROOF

We start by introducing the sufficient influence assumption of Theorem 1. Although it may seem intimidating at first, the reason why it is necessary will become clear when we prove the theorem.

Assumption 6 (Sufficient influence of $\mathbf{a}$ (nonparametric/continuous)) *For almost all $\mathbf{z} \in \mathbb{R}^{d_z}$ (i.e. except on a set with zero Lebesgue measure) and all $\ell \in [d_a]$, there exists*

$$\{(t_{(r)}, \tau_{(r)}, \mathbf{z}_{(r)}, \mathbf{a}_{(r)})\}_{r=1}^{|\text{Ch}_\ell^{\mathbf{a}}|},$$

such that $t_{(r)} \in [T]$, $\tau_{(r)} < t_{(r)}$, $\mathbf{z}_{(r)} \in \mathbb{R}^{d_z \times (t_{(r)}-1)}$, $\mathbf{a}_{(r)} \in \mathcal{A}^{t_{(r)}}$ and

$$\text{span} \left\{ H_{\mathbf{z},\mathbf{a}}^{t_{(r)},\tau_{(r)}} \log p(\mathbf{z} \mid \mathbf{z}_{(r)}, \mathbf{a}_{(r)})_{\cdot,\ell} \right\}_{r=1}^{|\text{Ch}_\ell^{\mathbf{a}}|} = \mathbb{R}_{\text{Ch}_\ell^{\mathbf{a}}}^{d_z}.$$

Proof [of Theorem 1] Recall equation (13), which we derived in Section 3.2:

$$\underbrace{H_{\mathbf{z},\mathbf{a}}^{t,\tau} \hat{q}(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t})}_{\subseteq \hat{\mathbf{G}}^{\mathbf{a}}} = D\mathbf{v}(\mathbf{z}^t)^\top \underbrace{H_{\mathbf{z},\mathbf{a}}^{t,\tau} q(\mathbf{v}(\mathbf{z}^t) \mid \mathbf{v}(\mathbf{z}^{<t}), \mathbf{a}^{<t})}_{\subseteq \mathbf{G}^{\mathbf{a}}}. \quad (16)$$

Notice that Assumption 6 holds only “almost everywhere”, i.e. on a set $\mathbb{R}^{d_z} \setminus E_0$ where E_0 has zero Lebesgue measure. Fix an arbitrary $\mathbf{z} \in \mathbb{R}^{d_z} \setminus E_0$. For notational convenience, define

$$\Lambda(\mathbf{z}, \gamma) := H_{\mathbf{z},\mathbf{a}}^{t,\tau} q(\mathbf{v}(\mathbf{z}) \mid \mathbf{v}(\mathbf{z}^{<t}), \mathbf{a}^{<t}) \quad \hat{\Lambda}(\mathbf{z}, \gamma) := H_{\mathbf{z},\mathbf{a}}^{t,\tau} \hat{q}(\mathbf{z} \mid \mathbf{z}^{<t}, \mathbf{a}^{<t}),$$

where $\gamma := (t, \tau, \mathbf{z}^{<t}, \mathbf{a}^{<t})$. This allows us to rewrite (16) with a much lighter notation:

$$\hat{\Lambda}(\mathbf{z}, \gamma) = D\mathbf{v}(\mathbf{z})^\top \Lambda(\mathbf{z}, \gamma). \quad (17)$$

Now, notice that the sufficient influence assumption (Assumption 6) requires that, for all $\ell \in [d_a]$ there exists $\{\gamma_{(r)}\}_{r=1}^{|\text{Ch}_\ell^{\mathbf{a}}|}$ such that $\text{span}\{\Lambda(\mathbf{z}, \gamma_{(r)})_{\cdot,\ell}\}_{r=1}^{|\text{Ch}_\ell^{\mathbf{a}}|} = \mathbb{R}_{\text{Ch}_\ell^{\mathbf{a}}}^{d_z}$. We can thus write

$$D\mathbf{v}(\mathbf{z})^\top \mathbb{R}_{\hat{\mathbf{G}}_{\cdot,\ell}^{\mathbf{a}}}^{d_z} = D\mathbf{v}(\mathbf{z})^\top \text{span}\{\Lambda(\mathbf{z}, \gamma_{(r)})_{\cdot,\ell}\}_{r=1}^{|\text{Ch}_\ell^{\mathbf{a}}|} = \text{span}\{\hat{\Lambda}(\mathbf{z}, \gamma_{(r)})_{\cdot,\ell}\}_{r=1}^{|\text{Ch}_\ell^{\mathbf{a}}|} \subseteq \mathbb{R}_{\hat{\mathbf{G}}_{\cdot,\ell}^{\mathbf{a}}}^{d_z} \quad (18)$$

Since $Dv(z)$ is invertible, there exists a permutation $P(z)$ such that $Dv(z)P(z)$ has no zero on its diagonal (Lemma 2). Let $C(z) := Dv(z)P(z)$. By left-multiplying (18) by $P(z)^\top$, we get

$$C(z)^\top \mathbb{R}_{\hat{G}_{\cdot,\ell}^a}^{d_z} \subseteq \mathbb{R}_{P(z)^\top \hat{G}_{\cdot,\ell}^a}^{d_z}. \quad (19)$$

We would like to show that $C(z)$ is G^a -preserving. Notice how the above equation is almost exactly the definition of G^a -preserving. All that is left to prove is that $P(z)^\top \hat{G}^a = G^a$.

We start by showing $P(z)^\top \hat{G}^a \supseteq G^a$. Take $(i, \ell) \in G^a$. Since $e_i \in \mathbb{R}_{\hat{G}_{\cdot,\ell}^a}^{d_z}$, equation (19) implies

$$C(z)^\top e_i = C(z)_{i,\cdot} \in \mathbb{R}_{P(z)^\top \hat{G}_{\cdot,\ell}^a}^{d_z}.$$

Since $C(z)_{i,i} \neq 0$ (all elements on its diagonal are nonzero), we must have that $(i, \ell) \in P(z)^\top \hat{G}^a$.

Now, since $\|P(z)^\top \hat{G}^a\|_0 = \|\hat{G}^a\|_0 \leq \|G^a\|_0$, we have $P(z)^\top \hat{G}^a = G^a$. This implies

$$C(z)^\top \mathbb{R}_{G^a}^{d_z \times d_a} \subseteq \mathbb{R}_{G^a}^{d_z \times d_a},$$

i.e. $C(z)$ is a G^a -preserving matrix, as desired.

To recap, we now have that, for all $z \in \mathbb{R}^{d_z} \setminus E_0$, there exists a permutation $P(z)$ s.t. $Dv(z)P(z)$ is G^a -preserving. We are not done yet, since, a priori, the permutation $P(z)$ can be different for different values of z , and we do not know what happens on the measure-zero set E_0 . What we need to show is that there exists a permutation P such that, for all z , $Dv(z)P$ is G^a -preserving. Lemma 12 in Appendix A.5 shows precisely this, by leveraging the continuity of $Dv(z)$ (v is a diffeomorphism and thus C^1).

Notice that $D(v \circ P)(z) = Dv(Pz)P$, which is G^a -preserving everywhere. Using Lemma 1, we conclude that the function $c := v \circ P$ is G^a -preserving. This concludes the proof. $\blacksquare$

Remark 5 (Alternative view on sufficient influence assumptions) *Assumption 6, and all sufficient influence assumptions we present later on, can be thought of in terms of linear independence of functions. By definition, a family of functions $(f^{(i)} : X \rightarrow \mathbb{R})_{i=1}^n$ is linearly independent when $\sum_i \alpha_i f^{(i)}(x) = 0$ for all $x \in X$ implies $\alpha_i = 0$ for all i . It turns out that Assumption 6 is equivalent to requiring that, for all $z \in \mathbb{R}^{d_z}$ and $\ell \in [d_a]$, the family of functions $(H_{z,a}^{t,\tau} \log p(z \mid z^{<t}, a^{<t})_{i,\ell})_{i \in \text{Ch}_\ell^a}$ (seen as functions of $t, \tau, z^{<t}$ and $a^{<t}$) is linearly independent. To see this, note that, in general, $(f^{(i)} : X \rightarrow \mathbb{R})_{i=1}^n$ is linearly independent iff there exist $x_1, \dots, x_n \in X$ s.t. the vectors $((f^{(1)}(x_i), \dots, f^{(n)}(x_i)))_{i=1}^n$ are linearly independent (see Appendix A.1 for a proof).*

3.8.2 SUFFICIENT INFLUENCE ASSUMPTION OF THEOREM 2 AND ITS PROOF

One can see that, if a is discrete, Theorem 1 cannot be applied because its sufficient influence assumption (Assumption 6) refers to the cross derivative of $\log p$ w.r.t. z^t and a^τ , which, of course, is not well defined when a is discrete. The discrete case is important to discuss interventions with unknown-targets as we did in Section 3.4.1, which is why we have a specialized result (Theorem 2) which has an analogous sufficient influence assumption based on *partial differences*.

Definition 16 (Partial difference) *Let us define*

$$\Delta_{a,\ell}^{\tau,\epsilon} D_z^t \log p(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t}) := D_z^t \log p(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t} + \epsilon \mathbf{E}^{(\ell,\tau)}) - D_z^t \log p(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t}),$$

where $\epsilon \in \mathbb{R}$ and $\mathbf{E}^{(\ell,\tau)}$ is a matrix with a one at entry (ℓ, τ) and zeros everywhere else.

One can see that $\Delta_{a,\ell}^{\tau,\epsilon} D_z^t \log p$ is essentially the discrete analog of $(H_{z,a}^{t,\tau} \log p)_{,\ell}$. Apart from this difference, the sufficient influence assumption for discrete $\mathbf{a}$ is the same as for continuous $\mathbf{a}$.

Assumption 7 (Sufficient influence of $\mathbf{a}$ (nonparametric/discrete)) *For almost all $\mathbf{z} \in \mathbb{R}^{d_z}$ (i.e. except on a set with zero Lebesgue measure) and all $\ell \in [d_a]$, there exists*

$$\{(t_{(r)}, \tau_{(r)}, \mathbf{z}_{(r)}, \mathbf{a}_{(r)}^{<t}, \epsilon_{(r)})\}_{r=1}^{|\text{Ch}_\ell^{\mathbf{a}}|},$$

such that $t_{(r)} \in [T]$, $\tau_{(r)} < t_{(r)}$, $\mathbf{z}_{(r)} \in \mathbb{R}^{d_z \times (t_{(r)}-1)}$, $\mathbf{a}_{(r)} \in \mathcal{A}^{t_{(r)}}$, $\epsilon_{(r)} \in \mathbb{R}$, $(\mathbf{a}_{(r)})_{,\tau_{(r)}} + \epsilon_{(r)} \mathbf{e}_\ell \in \mathcal{A}$ and

$$\text{span} \left\{ \Delta_{a,\ell}^{\tau_{(r)}, \epsilon_{(r)}} D_z^{t_{(r)}} \log p(\mathbf{z} \mid \mathbf{z}_{(r)}, \mathbf{a}_{(r)}) \right\}_{r=1}^{|\text{Ch}_\ell^{\mathbf{a}}|} = \mathbb{R}_{\text{Ch}_\ell^{\mathbf{a}}}^{d_z}.$$

We can now provide a proof of Theorem 2. Note that it is almost identical to the proof of Theorem 1 except for the very first steps where we take a partial difference instead of a partial derivative.

Proof [of Theorem 2] We recall equation (8) derived in Section 3.2:

$$D_z^t \hat{q}(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t}) = D_z^t q(\mathbf{v}(\mathbf{z}^t) \mid \mathbf{v}(\mathbf{z}^{<t}), \mathbf{a}^{<t}) D\mathbf{v}(\mathbf{z}^t) + \eta(\mathbf{z}^t) \in \mathbb{R}^{1 \times d_z}. \quad (20)$$

Now, instead of differentiating w.r.t. $\mathbf{a}_\ell^\tau$ for some $\tau < t$ and $\ell \in [d_a]$, we are going to take a partial difference. That is, we evaluate the above equation on at $\mathbf{a}^{<t}$ and $\mathbf{a}^{<t} + \epsilon \mathbf{E}^{(\ell,\tau)}$ and $\epsilon \in \mathbb{R}$, where $\mathbf{E}^{(\ell,\tau)}$ is a “one-hot matrix”, while keeping everything else constant, and take the difference. This yields:

$$\begin{aligned} & [D_z^t \hat{q}(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t} + \epsilon \mathbf{E}^{(\ell,\tau)}) - D_z^t \hat{q}(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t})]^\top \\ &= D\mathbf{v}(\mathbf{z}^t)^\top [D_z^t q(\mathbf{v}(\mathbf{z}^t) \mid \mathbf{v}(\mathbf{z}^{<t}), \mathbf{a}^{<t} + \epsilon \mathbf{E}^{(\ell,\tau)}) - D_z^t q(\mathbf{v}(\mathbf{z}^t) \mid \mathbf{v}(\mathbf{z}^{<t}), \mathbf{a}^{<t})]^\top \\ & \Delta_{a,\ell}^{\tau,\epsilon} D_z^t \hat{q}(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t})^\top = D\mathbf{v}(\mathbf{z}^t)^\top \Delta_{a,\ell}^{\tau,\epsilon} D_z^t q(\mathbf{v}(\mathbf{z}^t) \mid \mathbf{v}(\mathbf{z}^{<t}), \mathbf{a}^{<t})^\top, \end{aligned} \quad (21)$$

where we used the notation for partial differences introduced in Definition 16. Notice that the difference on the left is $\subseteq \hat{\mathbf{G}}_{\cdot,\ell}^{\mathbf{a}}$ and the difference on the right is $\subseteq \mathbf{G}_{\cdot,\ell}^{\mathbf{a}}$. This equation is thus analogous to (16) from the continuous case. For that reason, we can employ a completely analogous strategy. Hence, we define

$$\hat{\Lambda}(\mathbf{z}^t, \gamma)_{,\ell} := \Delta_{a,\ell}^{\tau,\epsilon} D_z^t \hat{q}(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t})^\top \quad \Lambda(\mathbf{z}^t, \gamma)_{,\ell} := \Delta_{a,\ell}^{\tau,\epsilon} D_z^t q(\mathbf{v}(\mathbf{z}^t) \mid \mathbf{v}(\mathbf{z}^{<t}), \mathbf{a}^{<t})^\top,$$

where $\gamma = (t, \tau, \mathbf{z}^{<t}, \mathbf{a}^{<t}, \epsilon)$. This notation allows us to rewrite (21) more compactly as

$$\underbrace{\hat{\Lambda}(\mathbf{z}^t, \gamma)}_{\subseteq \hat{\mathbf{G}}^{\mathbf{a}}} = \mathbf{L}(\mathbf{z}^t)^\top \underbrace{\Lambda(\mathbf{z}^t, \gamma)}_{\subseteq \mathbf{G}^{\mathbf{a}}}.$$

From here, the rest of the argument is exactly analogous to the proof of Theorem 1. ■

3.8.3 SUFFICIENT INFLUENCE ASSUMPTION OF THEOREM 3 AND ITS PROOF

We now introduce the sufficient influence assumption of Theorem 3, which showed how regularizing the temporal dependency graph $\hat{G}^z$ to be sparse can result in disentanglement. Again, it is very similar to other sufficient influence assumptions we have seen so far.

Assumption 8 (Sufficient influence of z (nonparametric)) *For almost all $z \in \mathbb{R}^{d_z}$ (i.e. except on a set with zero Lebesgue measure), there exists*

$$\{(t_{(r)}, \tau_{(r)}, \mathbf{z}_{(r)}, \mathbf{a}_{(r)})\}_{r=1}^{\|\mathbf{G}^z\|_0},$$

such that $t_{(r)} \in [T]$, $\tau_{(r)} < t_{(r)}$, $\mathbf{z}_{(r)} \in \mathbb{R}^{d_z \times (t_{(r)}-1)}$, $\mathbf{a}_{(r)} \in \mathcal{A}^{t_{(r)}}$, $\mathbf{z} = \mathbf{z}_{(r)}^{\tau_{(r)}}$ and

$$\text{span} \left\{ H_{z,z}^{t_{(r)}, \tau_{(r)}} q(\mathbf{z} \mid \mathbf{z}_{(r)}, \mathbf{a}_{(r)}) \right\}_{r=1}^{\|\mathbf{G}^z\|_0} = \mathbb{R}_{\hat{G}^z}^{d_z}.$$

Proof [of Theorem 3] We recall equation (15) derived in Section 3.2:

$$\underbrace{H_{z,z}^{t,\tau} \hat{q}(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t})}_{\subseteq \hat{G}^z} = D\mathbf{v}(\mathbf{z}^t)^\top \underbrace{H_{z,z}^{t,\tau} q(\mathbf{v}(\mathbf{z}^t) \mid \mathbf{v}(\mathbf{z}^{<t}), \mathbf{a}^{<t})}_{\subseteq \mathbf{G}^z} D\mathbf{v}(\mathbf{z}^\tau).$$

This equation holds for all pairs of $\mathbf{z}^t$ and $\mathbf{z}^\tau$ in $\mathbb{R}^{d_z}$. We can thus evaluate it at a point such that $\mathbf{z}^t = \mathbf{z}^\tau$, which yields

$$H_{z,z}^{t,\tau} \hat{q}(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t}) = D\mathbf{v}(\mathbf{z}^t)^\top H_{z,z}^{t,\tau} q(\mathbf{v}(\mathbf{z}^t) \mid \mathbf{v}(\mathbf{z}^{<t}), \mathbf{a}^{<t}) D\mathbf{v}(\mathbf{z}^t). \quad (22)$$

Recall that Assumption 8 holds for all $\mathbf{z}^t \in \mathbb{R}^{d_z} \setminus E_0$ where E_0 has Lebesgue measure zero. Fix an arbitrary $\mathbf{z}^t \in \mathbb{R}^{d_z} \setminus E_0$ and set $\mathbf{z}^\tau = \mathbf{z}^t$. Let us define

$$\Lambda(\mathbf{z}^t, \gamma) := H_{z,z}^{t,\tau} q(\mathbf{v}(\mathbf{z}^t) \mid \mathbf{v}(\mathbf{z}^{<t}), \mathbf{a}^{<t}) \quad \hat{\Lambda}(\mathbf{z}^t, \gamma) := H_{z,z}^{t,\tau} \hat{q}(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t}),$$

where $\gamma := (t, \tau, \mathbf{z}_{-t}^{<t}, \mathbf{a}^{<t})$ and $\mathbf{z}_{-t}^{<t}$ is $\mathbf{z}^{<t}$ but without $\mathbf{z}^\tau$. We can now rewrite (22) compactly as

$$\hat{\Lambda}(\mathbf{z}^t, \gamma) = D\mathbf{v}(\mathbf{z}^t)^\top \Lambda(\mathbf{z}^t, \gamma) D\mathbf{v}(\mathbf{z}^t).$$

Now, notice that the sufficient influence assumption (Assumption 8) requires that, there exists $\{\gamma_{(r)}\}_{r=1}^{\|\mathbf{G}^z\|_0}$ such that $\text{span}\{\Lambda(\mathbf{z}^t, \gamma_{(r)})\}_{r=1}^{\|\mathbf{G}^z\|_0} = \mathbb{R}_{\hat{G}^z}^{d_z \times d_z}$. We can thus write

$$\begin{aligned} D\mathbf{v}(\mathbf{z}^t)^\top \text{span}\{\Lambda(\mathbf{z}^t, \gamma_{(r)})\}_{r=1}^{\|\mathbf{G}^z\|_0} D\mathbf{v}(\mathbf{z}^t) &= \text{span}\{\hat{\Lambda}(\mathbf{z}^t, \gamma_{(r)})\}_{r=1}^{\|\mathbf{G}^z\|_0} \subseteq \mathbb{R}_{\hat{G}^z}^{d_z \times d_z} \\ \implies D\mathbf{v}(\mathbf{z}^t)^\top \mathbb{R}_{\hat{G}^z}^{d_z \times d_z} D\mathbf{v}(\mathbf{z}^t) &\subseteq \mathbb{R}_{\hat{G}^z}^{d_z \times d_z} \end{aligned} \quad (23)$$

Since $D\mathbf{v}(\mathbf{z})$ is invertible, there exists a permutation $\mathbf{P}(\mathbf{z})$ such that $D\mathbf{v}(\mathbf{z})\mathbf{P}(\mathbf{z})$ has no zero on its diagonal (Lemma 2). Let $\mathbf{C}(\mathbf{z}) := D\mathbf{v}(\mathbf{z})\mathbf{P}(\mathbf{z})$. If we left and right-multiply (23) by $\mathbf{P}(\mathbf{z})^\top$ and $\mathbf{P}(\mathbf{z})$, respectively, we obtain

$$\mathbf{C}(\mathbf{z}^t)^\top \mathbb{R}_{\hat{G}^z}^{d_z \times d_z} \mathbf{C}(\mathbf{z}^t) \subseteq \mathbb{R}_{\mathbf{P}(\mathbf{z})^\top \hat{G}^z \mathbf{P}(\mathbf{z})}^{d_z \times d_z}. \quad (24)$$

We now show that $G^z \subseteq P(z)^\top \hat{G}^z P(z)$. Take $(i, j) \in G^z$. Since $e_i e_j^\top \in \mathbb{R}_{G^z}^{d_z \times d_z}$, equation (24) implies

$$C(z^t)^\top e_i e_j^\top C(z^t) = (C(z^t)_{i,\cdot})^\top C(z^t)_{j,\cdot} \subseteq \mathbb{R}_{P(z)^\top \hat{G}^z P(z)}^{d_z \times d_z} \quad (25)$$

Since $C(z^t)_{i,i} C(z^t)_{j,j} \neq 0$ (recall the diagonal of $C(z^t)$ has no zero), we must have $(i, j) \in P(z)^\top \hat{G}^z P(z)$. This shows that $G^z \subseteq P(z)^\top \hat{G}^z P(z)$.

Since $\|P(z)^\top \hat{G}^z P(z)\|_0 = \|\hat{G}^z\|_0 \leq \|G^z\|_0$, we must have $G^z = P(z)^\top \hat{G}^z P(z)$, which yields

$$C(z^t)^\top \mathbb{R}_{G^z}^{d_z \times d_z} C(z^t) \subseteq \mathbb{R}_{G^z}^{d_z \times d_z}. \quad (26)$$

We are now going to show that the above implies that $C(z^t)$ is both G^z -preserving and $(G^z)^\top$ -preserving. Start by rewriting (25) as follows:

$$\text{for all } (i, j) \in G^z, (C(z^t)_{i,\cdot})^\top C(z^t)_{j,\cdot} \subseteq \mathbb{R}_{G^z}^{d_z \times d_z}. \quad (27)$$

We start by showing G^z -preservation. To do so, we leverage the characterization of Proposition 3. We must show that $G_{i,\cdot}^z \not\subseteq G_{j,\cdot}^z$ implies $C(z^t)_{i,j} = 0$. Because $G_{i,\cdot}^z \not\subseteq G_{j,\cdot}^z$, there must exist k s.t. $G_{i,k}^z = 1$ and $G_{j,k}^z = 0$. We thus have, by (27), that $(C(z^t)_{i,\cdot})^\top C(z^t)_{k,\cdot} \subseteq \mathbb{R}_{G^z}^{d_z \times d_z}$. Because $G_{j,k}^z = 0$, we have $C(z^t)_{i,j} C(z^t)_{k,k} = 0$. But since $C(z^t)_{k,k} \neq 0$, we must have that $C(z^t)_{i,j} = 0$, as desired. To show $(G^z)^\top$ -preservation, one can use a completely analogous argument.

We showed that $C(z^t)$ is G^z -preserving and $(G^z)^\top$ -preserving. It is easy to verify that this is equivalent to being $[G^z (G^z)^\top]$ -preserving (where $[\cdot]$ stands for column concatenation). This remark will be useful below.

Similarly to the proof of Theorem 1, we must now show that there exists a single permutation that works for all $z^t \in \mathbb{R}^{d_z}$. To achieve this, we use Lemma 12 with $G := [G^z (G^z)^\top]$ and $L(z) := Dv(z)$. This allows us to say that there exists a permutation P such that $Dv(z)P$ is $[G^z (G^z)^\top]$ -preserving for all z (not “almost all”).

Notice that $D(v \circ P)(z) = Dv(Pz)P$, which is $[G^z (G^z)^\top]$ -preserving everywhere. Using Lemma 1, we conclude that the function $c := v \circ P$ is $[G^z (G^z)^\top]$ -preserving. $\blacksquare$

3.9 Examples to Illustrate the Scope of the Theory

In this section, we provide several examples to gain better intuition as to when our results apply. Specifically, we will provide mathematically concrete examples of latent models $p(z^t \mid z^{<t}, a^{<t})$ illustrating the various sufficient influence assumptions we introduced. All these examples are summarized in Table 2.

Even though our results are nonparametric, we will concentrate on the special case of Gaussian models which are useful to get a good intuition of what the sufficient influence assumptions mean. The following simple lemma will be useful in the following examples. We present it without proof, as it can be derived from simple computations.

Lemma 3 *Let $p(z) = \mathcal{N}(z; \mu, \Sigma)$ where $\mu \in \mathbb{R}^{d_z}$ and $\Sigma := \text{diag}(\sigma_1^2, \dots, \sigma_{d_z}^2)$. Then*

$$D_z \log p(z) = -[(z_1 - \mu_1)/\sigma_1^2, \dots, (z_{d_z} - \mu_{d_z})/\sigma_{d_z}^2] \in \mathbb{R}^{1 \times d_z}.$$

3.9.1 CONTINUOUS AUXILIARY VARIABLES (THEOREM 1)

We start by illustrating Assumption 6 from Theorem 1. Example 9 assumes that we observe continuous actions that target each latent factor individually, while Example 10 gives a multi-target example.

Example 9 (Sufficient influence for continuous single-target actions) We make Example 4 more concrete by explicitly specifying a latent transition model. Recall the situation depicted in Figure 1 where $\mathbf{z}_1$ is the tree position, $\mathbf{z}_2$ is the robot position, and $\mathbf{z}_3$ is the ball position ($d_z = 3$). Assume $\mathbf{a} \in [-1, 1]$ corresponds to the amount of torque applied to the wheels of the robot. We thus have $\mathbf{G}^a = [0, 1, 0]^\top$, i.e., $\mathbf{a}$ affects only the robot position $\mathbf{z}_2$. For this example, $\mathbf{G}^z$ can be anything. Let $p(\mathbf{z}^t | \mathbf{z}^{t-1}, \mathbf{a}) = \mathcal{N}(\mathbf{z}^t; \boldsymbol{\mu}(\mathbf{z}^{t-1}, \mathbf{a}), \sigma^2 \mathbf{I})$ where

$$\boldsymbol{\mu}(\mathbf{z}^{t-1}, \mathbf{a}) := \mathbf{z}^{t-1} + \mathbf{g}(\mathbf{z}^{t-1}) + \mathbf{a} \cdot \mathbf{G}^a.$$

where $\mathbf{g} : \mathbb{R}^{d_z} \rightarrow \mathbb{R}^{d_z}$ is some function that satisfies the dependency graph $\mathbf{G}^z$ (e.g. $\mathbf{g}(\mathbf{z}) := \mathbf{W}\mathbf{z}$ where $\mathbf{W} \in \mathbb{R}^{d_z \times d_z}$). If no torque is applied ($\mathbf{a} = 0$), then the position of the robots is determined by the dynamics of the system. However, adding a positive or negative torque ($\mathbf{a} \neq 0$) nudges the robot to the right or to the left. Using Lemma 3, we can compute

$$H_{z,a}^t \log p(\mathbf{z}^t | \mathbf{z}^{t-1}, \mathbf{a}) = [0, 1/\sigma^2, 0]^\top,$$

which spans $\mathbb{R}_{\{2\}}^3$ and thus Assumption 6 holds.

Example 10 (Sufficient influence for continuous multi-target actions) We make Example 5 more concrete by specifying an explicit latent model. Recall that $\mathbf{G}^a$ is given by Figure 3c with $d_z = d_a = 3$. Assume there are no temporal dependencies ($T = 1$), that $\mathbf{a} \in \mathbb{R}^3$ and that the latent model is given by $p(\mathbf{z} | \mathbf{a}) = \mathcal{N}(\mathbf{z}; \boldsymbol{\mu}(\mathbf{a}), \sigma^2 \mathbf{I})$ where

$$\boldsymbol{\mu}(\mathbf{a}) := \begin{bmatrix} \mathbf{a}_1 \\ \mathbf{a}_1^2 \\ 0 \end{bmatrix} + \begin{bmatrix} \mathbf{a}_2 \\ 0 \\ \mathbf{a}_2^2 \end{bmatrix} + \begin{bmatrix} 0 \\ \mathbf{a}_3 \\ \mathbf{a}_3^2 \end{bmatrix}. \quad (28)$$

Using Lemma 3, we can compute

$$H_{z,a} q(\mathbf{z} | \mathbf{a}) = \frac{1}{\sigma^2} \begin{bmatrix} 1 & 1 & 0 \\ 2\mathbf{a}_1 & 0 & 1 \\ 0 & 2\mathbf{a}_2 & 2\mathbf{a}_3 \end{bmatrix}.$$

Consider $\ell = 1$ so that $\mathbf{Ch}_1^a = \{1, 2\}$. We can see that $H_{z,a} q(\mathbf{z} | \mathbf{a} = 0)_{\cdot,1} = [1, 0, 0]^\top$ and $H_{z,a} q(\mathbf{z} | \mathbf{a} = \mathbf{e}_1)_{\cdot,1} = [1, 2, 0]^\top$ span $\mathbb{R}_{\{1,2\}}^3$. Analogous conclusions can also be reached for $\ell = 2, 3$, which shows that Assumption 6 holds.

Now suppose that instead $\boldsymbol{\mu}(\mathbf{a})$ is a linear map, i.e., $\boldsymbol{\mu}(\mathbf{a}) := \mathbf{W}\mathbf{a}$ where $\mathbf{W} \in \mathbb{R}_{\mathbf{G}^a}^{d_z \times d_a}$. This would imply that $H_{z,a} q(\mathbf{z} | \mathbf{a}) \propto \mathbf{W}$, which means it cannot satisfy the sufficient influence assumption (unless $\|\mathbf{G}_{\cdot,\ell}^a\|_0 \leq 1$ for all ℓ).

3.9.2 DISCRETE AUXILIARY VARIABLES OR INTERVENTIONS (THEOREM 2)

We now provide three concrete examples of latent models $p(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t})$ that satisfy Assumption 7, from Theorem 2. Here, we interpret the discrete auxiliary variable $\mathbf{a}$ as an *intervention index*, as discussed in Section 3.4.1, but note that other interpretations are possible (such as $\mathbf{a}$ as an action). Recall that our identifiability result does not require the knowledge of the targets of the interventions, these can be learned.

Example 11 shows how single target interventions can be used to obtain complete disentanglement without temporal dependencies, Example 12 shows how multi-target interventions can be leveraged for disentanglement if temporal dependencies are present, and Example 13 shows how *grouped* multi-target interventions allow disentanglement even when there are no time dependencies (Remark 6).

Example 11 (Single-target interventions for complete disentanglement without time) *We make Example 3 more concrete by specifying an explicit latent model. Assume $d_a = d_z$ and that $\mathbf{a} \in \mathcal{A} := \{\mathbf{0}, \mathbf{e}_1, \dots, \mathbf{e}_{d_a}\}$ is interpreted as an intervention index (see Section 3.4.1). Furthermore, Example 3 assumed $\mathbf{G}^a = \mathbf{I}$, i.e. each latent factor is targeted once by an intervention that targets only this factor (the example actually allowed to add arbitrary columns to $\mathbf{G}^a$, i.e. adding more interventions, without compromising complete disentanglement). Assume there are no temporal dependencies, i.e., $T = 1$, and that $p(\mathbf{z} \mid \mathbf{a}) := \mathcal{N}(\mathbf{z}; \boldsymbol{\mu}(\mathbf{a}), \text{diag}(\boldsymbol{\sigma}^2(\mathbf{a})))$ with*

$$\boldsymbol{\mu}(\mathbf{a}) := \boldsymbol{\mu} \odot \mathbf{a} \quad \text{and} \quad \boldsymbol{\sigma}^2(\mathbf{a}) := \mathbb{1} + \boldsymbol{\delta} \odot \mathbf{a},$$

where $\odot$ denotes the Hadamard product (a.k.a. the element-wise product), $\boldsymbol{\mu} \in \mathbb{R}^{d_z}$ is the vector of means for each intervention and $\boldsymbol{\delta} \in \mathbb{R}^{d_z}$ is the vector of shifts in variance for all interventions. Thus, in the observational setting ($\mathbf{a} = \mathbf{0}$), we have $\boldsymbol{\mu}(\mathbf{a}) = \mathbf{0}$ and $\boldsymbol{\sigma}(\mathbf{a}) = \mathbb{1}$ while in the ℓ th intervention ($\mathbf{a} = \mathbf{e}_\ell$), the mean and variance of the targeted latent shift are the same, i.e. $\boldsymbol{\mu}(\mathbf{a}) = \boldsymbol{\mu}_\ell \mathbf{e}_\ell$ and $\boldsymbol{\sigma}^2(\mathbf{a}) = \mathbb{1} + \boldsymbol{\delta}_\ell \mathbf{e}_\ell$ (assuming that the shifted variance is > 0). Using Lemma 3, we can compute

$$\Delta_{a,\ell}^{\epsilon=1} D_z \log p(\mathbf{z} \mid \mathbf{a} = \mathbf{0}) := D_z \log p(\mathbf{z} \mid \mathbf{a} = \mathbf{e}_\ell) - D_z \log p(\mathbf{z} \mid \mathbf{a} = \mathbf{0}) = \frac{\boldsymbol{\mu}_\ell + \boldsymbol{\delta}_\ell \mathbf{z}_\ell}{1 + \boldsymbol{\delta}_\ell} \mathbf{e}_\ell,$$

which must span $\mathbb{R}_{\{\ell\}}^{d_z}$ unless $\boldsymbol{\mu}_\ell + \boldsymbol{\delta}_\ell \mathbf{z}_\ell = \mathbf{0}$. However, note that when, for all ℓ , $\boldsymbol{\mu}_\ell \neq \mathbf{0}$ or $\boldsymbol{\delta}_\ell \neq \mathbf{0}$ (i.e. all interventions truly have an effect), the set $\{\mathbf{z} \in \mathbb{R}^{d_z} \mid \boldsymbol{\mu}_\ell + \boldsymbol{\delta}_\ell \mathbf{z}_\ell = \mathbf{0} \text{ for some } \ell\}$ has zero Lebesgue measure in $\mathbb{R}^{d_z}$, which is allowed by Assumption 7.

Remark 6 (Potential issues with multi-target interventions without time) *What if an intervention targets more than one latent at a time? Can it still satisfy the sufficient influence assumption? We will now see that, without time-dependencies ($T = 1$), it is impossible. Consider the simple situation where $d_z = 3$, $d_a = 1$, $\mathbf{a} \in \{0, 1\}$ and $\mathbf{G}^a = [1, 1, 0]^\top$, i.e., there is a single intervention targeting $\mathbf{z}_1$ and $\mathbf{z}_2$. In that case, there is a single possible difference vector, which is*

$$\Delta_a^{\epsilon=1} D_z \log p(\mathbf{z} \mid \mathbf{a} = 0) = D_z \log p(\mathbf{z} \mid \mathbf{a} = 1) - D_z \log p(\mathbf{z} \mid \mathbf{a} = 0) \in \mathbb{R}_{\{1,2\}}^{d_z}.$$

Since this is the only difference vector, we can see that we cannot span the 2-dimensional space $\mathbb{R}_{\{1,2\}}^{d_z}$. Therefore, to leverage multi-target interventions in our framework, more “variability” is required. Example 12 below shows how temporal dependencies can provide this additional variability, while Example 13 shows how having “groups” of interventions known to have the same (unknown) targets can also provide the required variability.

Example 12 (Multi-target interventions for complete disentanglement with time) We make Example 5 more concrete by specifying an explicit latent model that satisfies Assumption 7. Recall $d_z = 3$, $d_a = 3$ and $\mathbf{G}^a$ is depicted in Figure 3c. This time, we assume there are temporal dependencies, i.e. $T > 1$ and $\mathbf{G}^z$ is non-trivial. Suppose $\mathbf{a} \in \mathcal{A} := \{\mathbf{0}, \mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3\}$ where $\mathbf{e}_\ell$ is the ℓ th one-hot and we interpret $\mathbf{0}$ to correspond to the observational setting and $\mathbf{e}_\ell$ to correspond to the ℓ th intervention. Recall that in this interpretation $\mathbf{G}^a$ describes which latent variable is targeted by each intervention. Let $p(\mathbf{z}^t \mid \mathbf{z}^{t-1}, \mathbf{a}) = \mathcal{N}(\mathbf{z}^t; \boldsymbol{\mu}(\mathbf{z}^{t-1}, \mathbf{a}), \sigma^2 \mathbf{I})$ where

$$\boldsymbol{\mu}(\mathbf{z}^{t-1}, \mathbf{a}) := \mathbf{z}^{t-1} + (\mathbb{I} - \mathbf{G}^a \mathbf{a}) \odot \mathbf{g}(\mathbf{z}^{t-1}),$$

where $\mathbf{g} : \mathbb{R}^{d_z} \rightarrow \mathbb{R}^{d_z}$ is some function that respects the graph $\mathbf{G}^z$ (e.g. $\mathbf{g}(\mathbf{z}) = \mathbf{W}\mathbf{z}$ where $\mathbf{W} \in \mathbb{R}_{\mathbf{G}^z}^{d_z \times d_z}$). The observational dynamics is then $\boldsymbol{\mu}(\mathbf{z}^t, \mathbf{a} = \mathbf{0}) = \mathbf{z}^{t-1} + \mathbf{g}(\mathbf{z}^{t-1})$ and the interventional settings correspond to zeroing out the elements of $\mathbf{g}(\mathbf{z}^{t-1})$ targeted by the intervention. Using Lemma 3, we can compute

$$\begin{aligned} \Delta_{\mathbf{a}, \ell}^{\epsilon=1} D_z q(\mathbf{z}^t \mid \mathbf{z}^{t-1}, \mathbf{a} = \mathbf{0}) \\ = D_z q(\mathbf{z}^t \mid \mathbf{z}^{t-1}, \mathbf{a} = \mathbf{e}_\ell) - D_z q(\mathbf{z}^t \mid \mathbf{z}^{t-1}, \mathbf{a} = \mathbf{0}) = -\frac{1}{\sigma^2} \mathbf{G}_{\cdot, \ell}^a \odot \mathbf{g}(\mathbf{z}^{t-1}). \end{aligned}$$

One can see that, as soon as the image of $\mathbf{g}$ spans $\mathbb{R}^{d_z}$, Assumption 7 is satisfied since we can choose values $\mathbf{z}_{(1)}, \dots, \mathbf{z}_{(d_z)} \in \mathbb{R}^{d_z}$ such that $\text{span}\{\mathbf{g}(\mathbf{z}_{(1)}), \dots, \mathbf{g}(\mathbf{z}_{(d_z)})\} = \mathbb{R}^{d_z}$, which implies $\text{span}\{\mathbf{G}_{\cdot, \ell}^a \odot \mathbf{g}(\mathbf{z}_{(1)}), \dots, \mathbf{G}_{\cdot, \ell}^a \odot \mathbf{g}(\mathbf{z}_{(d_z)})\} = \mathbb{R}_{\text{Ch}_\ell^a}^{d_z}$. An example of a transition function $\mathbf{g}$ that satisfies this property is $\mathbf{g}(\mathbf{z}) := \mathbf{W}\mathbf{z}$ where $\mathbf{W} \in \mathbb{R}_{\mathbf{G}^z}^{d_z \times d_z}$ is invertible.

Note that even if the temporal dependencies are not sparse, they are still helpful for identifiability as they make it more likely to satisfy the sufficient influence assumption (Assumption 7).

Example 13 (Grouped multi-target interventions for disentanglement without time) In this example, we assume that there are no temporal dependencies ($T = 1$) and that the learner has access to d_a groups of interventions where the interventions belonging to the ℓ th group are known to target the same latent variables given by $\mathbf{G}_{\cdot, \ell}^a$ (these targets are unknown). Here is how our framework can accommodate this setting: Given that we have d_a groups of interventions where the ℓ th group contains k_ℓ interventions, we set $\mathcal{A} := \{\mathbf{0}, 1\mathbf{e}_1, \dots, k_1\mathbf{e}_1, 1\mathbf{e}_2, \dots, k_2\mathbf{e}_2, \dots, k_{d_a}\mathbf{e}_{d_a}\}$. In this setting, $\mathbf{a} = j\mathbf{e}_\ell$ corresponds to the j th intervention of the ℓ th group. Moreover, the sufficient influence assumption requires that interventions within a group ℓ span $\mathbb{R}_{\text{Ch}_\ell^a}^{d_z}$. More precisely, we need $\text{span}\{\Delta_{\mathbf{a}, \ell}^\epsilon D_z \log p(\mathbf{z} \mid \mathbf{a} = \mathbf{0})\}_{\epsilon=1}^{k_\ell} = \mathbb{R}_{\text{Ch}_\ell^a}^{d_z}$.

3.9.3 TEMPORAL DEPENDENCIES (THEOREM 3)

Finally, we provide an example (Example 14) where temporal dependencies alone (no auxiliary variable $\mathbf{a}$) are enough to disentangle. We start with an important remark about the sufficient influence assumption of Theorem 3.

Remark 7 (Auxiliary variables or non-Markovianity are required) An important observation is that, if the transition model does not have an auxiliary variable $\mathbf{a}$ and is Markovian, i.e., $p(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t}) = p(\mathbf{z}^t \mid \mathbf{z}^{t-1})$, then Assumption 8 cannot be satisfied (except in trivial circumstances). To see this, simply note that, in that case, $H_{z, z}^{t, t-1} q(\mathbf{z}^t \mid \mathbf{z}^{t-1})$ depends only on $\mathbf{z}^{t-1}$, which is

forced to be equal to $\mathbf{z}^t$. This means that the span of the Hessian must be at most one-dimensional, which means that the assumption cannot hold as soon as $\|\mathbf{G}^z\|_0 > 1$. Therefore, when no auxiliary variable $\mathbf{a}$ is observed, Assumption 8 requires that the transition model be non-Markovian. In Example 14, we provide a concrete example of a transition model without auxiliary variable $\mathbf{a}$ that satisfies this assumption. We will also see in Section 4 that if the transition model $p(\mathbf{z}^t | \mathbf{z}^{t-1})$ is in the exponential family, this assumption can be relaxed so that non-Markovianity is no longer required.

Example 14 (Sparse temporal dependencies for disentanglement without auxiliary variables)

We continue with Examples 6 & 7 which were based on the graphs $\mathbf{G}^z$ depicted in Figures 3d & 3e, respectively. Assume that no action is observed, i.e., we can only take advantage of the sparsity of $\mathbf{G}^z$ to disentangle. Examples 6 & 7 already showed that these graph structures allow complete disentanglement, as long as the sufficient influence assumption for $\mathbf{z}$ (Assumption 8) is satisfied. We now provide concrete transition models $p(\mathbf{z}^t | \mathbf{z}^{<t})$ that satisfy this requirement. Similarly to previous examples, assume $p(\mathbf{z}^t | \mathbf{z}^{<t}) = \mathcal{N}(\mathbf{z}^t | \boldsymbol{\mu}(\mathbf{z}^{t-1}, \mathbf{z}^{t-2}), \sigma^2 \mathbf{I})$ where

$$\boldsymbol{\mu}(\mathbf{z}^{t-1}, \mathbf{z}^{t-2}) := \mathbf{z}^{t-1} + \mathbf{W}(\mathbf{z}^{t-2})\mathbf{z}^{t-1},$$

where $\mathbf{W} : \mathbb{R}^{d_z} \rightarrow \mathbb{R}^{d_z \times d_z}$ is some function of $\mathbf{z}^{t-2}$. Using Lemma 3, we can derive

$$H_{z,z}^{t,t-1} q(\mathbf{z}^t | \mathbf{z}^{<t}) = \frac{1}{\sigma^2} [\mathbf{I} + \mathbf{W}(\mathbf{z}^{t-2})].$$

Thus, Assumption 8 holds when there exists $\{\mathbf{z}_{(r)}^{t-2}\}_{r=1}^{\|\mathbf{G}^z\|_0}$ such that

$$\text{span}\{\mathbf{I} + \mathbf{W}(\mathbf{z}_{(r)}^{t-2})\}_{r=1}^{\|\mathbf{G}^z\|_0} = \mathbb{R}^{d_z \times d_z}. \quad (29)$$

One can directly see that, if $\mathbf{W}(\mathbf{z}^{t-2})$ was actually constant in $\mathbf{z}^{t-2}$, the assumption could not hold (unless $\|\mathbf{G}^z\|_0 \leq 1$). This case would correspond to a simple linear model of the form $\boldsymbol{\mu}(\mathbf{z}^{t-1}) := \mathbf{z}^{t-1} + \mathbf{W}\mathbf{z}^{t-1}$. Our theory suggests that this transition function is “too simple” to allow disentanglement.

Nevertheless, we can find examples satisfying (29). For example, if $\mathbf{G}^z = \mathbf{I}$, we can take

$$\mathbf{W}(\mathbf{z}) = \begin{bmatrix} z_1 & 0 & 0 \\ 0 & z_2 & 0 \\ 0 & 0 & z_3 \end{bmatrix}$$

and see that the family of functions $(1 + z_1, 1 + z_2, 1 + z_3)$ is linearly independent (when seen as functions from $\mathbb{R}^3$ to $\mathbb{R}$). By Lemma 5 in the appendix, this is equivalent to the existence of $\mathbf{z}_{(1)}, \mathbf{z}_{(2)}, \mathbf{z}_{(3)} \in \mathbb{R}^{d_z}$ such that (29) holds (see also Remark 5). In other words, the sufficient influence assumption holds. In the case where $\mathbf{G}^z$ is lower triangular like in Figure 3e, one can take

$$\mathbf{W}(\mathbf{z}) = \begin{bmatrix} z_1 & 0 & 0 \\ z_2^2 & z_2 & 0 \\ z_3^3 & z_3^2 & z_3 \end{bmatrix}$$

and see that the family of functions $(1 + z_1, 1 + z_2, 1 + z_3, z_2^2, z_3^2, z_3^3)$ is linearly independent, which similarly implies the existence of $\mathbf{z}_{(1)}, \dots, \mathbf{z}_{(6)} \in \mathbb{R}^{d_z}$ such that (29) holds.

Example 16 will show how one can leverage the exponential family assumption to allow Markovianity even without auxiliary variables.

4 Partial Disentanglement via Mechanism Sparsity in Exponential Families

The goal of this section is to understand how restricting the transition model to be in the *exponential family* allows us to weaken the sufficient influence assumption of Theorem 3. Section 4.1 introduces the exponential family assumption. Section 4.2 follows Khemakhem et al. (2020a) and shows that this additional assumption guarantees that the entanglement map v is “quasi-linear”, which means $v(z) := s^{-1}(Ls(z) + b)$, where L is a matrix and s is an element-wise invertible function. Section 4.3 will introduce an identifiability result analogous to Theorem 3 for sparse $\hat{G}^z$ that takes advantage of the quasi-linearity of v to weaken Assumption 8 (sufficient influence of z). We also briefly discuss an additional result from Appendix B.4 that shows connections between the non-parametric sufficient influence assumptions of this work (Assumptions 7 & 8) and their counterparts in Lachapelle et al. (2022) (Assumptions 11 & 12).

4.1 Exponential Family Latent Transition Models

We will assume that the conditional densities $p(z_i^t | z^{<t}, a^{<t})$ are from an *exponential family* (Wainwright and Jordan, 2008):

Assumption 9 (Exponential family transition model) For all $i \in [d_z]$, we have

$$p(z_i^t | z^{<t}, a^{<t}) = h_i(z_i^t) \exp\{s_i(z_i^t)^\top \lambda_i(z^{<t}, a^{<t}) - \psi_i(z^{<t}, a^{<t})\}. \quad (30)$$

Well-known distributions which belong to this family include the Gaussian and beta distributions. In the Gaussian case, the *sufficient statistic* is $s_i(z) := (z, z^2)$ and the *base measure* is $h_i(z) := \frac{1}{\sqrt{2\pi}}$. The function $\lambda_i(z^{<t}, a^{<t})$ outputs the *vectors of natural parameters* for the conditional distribution and can be itself parametrized, for instance, by a multi-layer perceptron (MLP) or a recurrent neural network (RNN). We will refer to the functions λ_i as the *mechanisms* or the *transition functions*. In the Gaussian case, the natural parameter is two-dimensional and is related to the usual parameters μ and σ^2 via the equation $(\lambda_1, \lambda_2) = (\frac{\mu}{\sigma^2}, -\frac{1}{2\sigma^2})$. We will denote by k the dimensionality of the natural parameter and that of the sufficient statistic (which are equal). Thus, $k = 2$ in the Gaussian case. The remaining term $\psi_i(z^{<t}, a^{<t})$ acts as a normalization constant.

We define $\lambda(z^{<t}, a^{<t}) \in \mathbb{R}^{kd_z}$ as the concatenation of all $\lambda_i(z^{<t}, a^{<t})$ and similarly for $s(z^t) \in \mathbb{R}^{kd_z}$. Similarly to the nonparameteric case, the learnable parameters are $\theta := (f, \lambda, G)$. Note that throughout, we assume that the sufficient statistic s is not learned and known in advance. With this notation, we can write the full transition model as

$$p(z^t | z^{<t}, a^{<t}) = h(z^t) \exp\{s(z^t)^\top \lambda(z^{<t}, a^{<t}) - \psi(z^{<t}, a^{<t})\},$$

where $h := \prod_{i=1}^{d_z} h_i$ and $\psi = \sum_{i=1}^{d_z} \psi_i$.

Remark 8 (Applying nonparametric identifiability results to exponential families) One can apply the nonparametric results (Theorems 1, 2 & 3) to models satisfying the exponential family assumption. In fact, all examples of Section 3.9 were Gaussians and thus are in the exponential family.

4.2 Conditions for Quasi-Linear Identifiability

In this section, we follow Khemakhem et al. (2020a) and show that the exponential family assumption combined with an additional sufficient variability assumption allows us to go from identifiability up to diffeomorphism (Definition 5) to identifiability up to quasi-linearity, which we define next:

Definition 17 (Quasi-linear equivalence) We say that two models $\theta := (\mathbf{f}, \boldsymbol{\lambda}, \mathbf{G})$ and $\tilde{\theta} := (\tilde{\mathbf{f}}, \tilde{\boldsymbol{\lambda}}, \tilde{\mathbf{G}})$ satisfying Assumptions 1, 2 & 9 are **equivalent up to quasi-linearity**, denoted $\theta \sim_{\text{lin}} \tilde{\theta}$, if and only if $\theta \sim_{\text{diff}} \tilde{\theta}$ and there exist an invertible matrix $\mathbf{L} \in \mathbb{R}^{kd_z \times kd_z}$ and a vector $\mathbf{b} \in \mathbb{R}^{kd_z}$ such that the map $\mathbf{v} := \mathbf{f}^{-1} \circ \tilde{\mathbf{f}}$ satisfies

$$\mathbf{s}(\mathbf{v}(\mathbf{z})) = \mathbf{L}\mathbf{s}(\mathbf{z}) + \mathbf{b}, \forall \mathbf{z} \in \mathbb{R}^{d_z}.$$

If the sufficient statistic $\mathbf{s}$ is invertible, one obtains

$$\mathbf{v}(\mathbf{z}) = \mathbf{s}^{-1}(\mathbf{L}\mathbf{s}(\mathbf{z}) + \mathbf{b}), \forall \mathbf{z} \in \mathbb{R}^{d_z}. \quad (31)$$

Equation (31) is particularly interesting, as it says that the mapping relating both representations is “almost” linear in the following sense: although the map is not necessarily linear because the sufficient statistic $\mathbf{s}$ might not be, the “mixing” between components is linear. Indeed, notice that the sufficient statistic $\mathbf{s}$ and its inverse operate “element-wise”. The mixing between components is only due to the matrix $\mathbf{L}$. This specific form simplifies a few steps in the identifiability proof, which might explain the popularity of this assumption in the literature on nonlinear ICA (Hyvärinen and Morioka, 2016; Khemakhem et al., 2020a,b; Hälvä and Hyvärinen, 2020; Morioka et al., 2021; Yang et al., 2021; Lachapelle et al., 2022; Liu et al., 2023; Xi and Bloem-Reddy, 2023).

The following theorem provides conditions to guarantee identifiability up to quasi-linearity. This is an adaptation and minor extension of Theorem 1 from Khemakhem et al. (2020a). For completeness, we provide a proof in Appendix B.2.

Theorem 4 (Conditions for linear identifiability - Adapted from Khemakhem et al. (2020a)) Let $\theta := (\mathbf{f}, \boldsymbol{\lambda}, \mathbf{G})$ and $\hat{\theta} := (\hat{\mathbf{f}}, \hat{\boldsymbol{\lambda}}, \hat{\mathbf{G}})$ be two models satisfying Assumptions 1, 2 & 9. Further assume that

1. **[Observational equivalence]** $\theta \sim_{\text{obs}} \hat{\theta}$ (Definition 4);
2. **[Minimal sufficient statistics]** For all i , the sufficient statistic $\mathbf{s}_i$ is minimal (see below).
3. **[Sufficient variability]** The natural parameter $\boldsymbol{\lambda}$ varies “sufficiently” as formalized by Assumption 10 (see below).

Then, $\theta \sim_{\text{lin}} \hat{\theta}$ (Def. 17).

The “minimal sufficient statistics” assumption is a standard one saying that $\mathbf{s}_i$ is defined appropriately to ensure that the parameters of the exponential family are identifiable (see e.g. Wainwright and Jordan (2008, p. 40)). See Definition 20 for a formal definition of minimality. The last assumption is sometimes called the *assumption of variability* (Hyvärinen et al., 2019), and requires that the conditional distribution of $\mathbf{z}^t$ depends “sufficiently strongly” on $\mathbf{z}^{<t}$ and/or $\mathbf{a}^{<t}$. We stress the fact that this assumption concerns the ground-truth data generating model θ .

Assumption 10 (Sufficient variability in exponential families) There exist $(\mathbf{z}_{(r)}, \mathbf{a}_{(r)})_{r=0}^{kd_z}$ in their respective supports such that the kd_z -dimensional vectors $(\boldsymbol{\lambda}(\mathbf{z}_{(r)}, \mathbf{a}_{(r)}) - \boldsymbol{\lambda}(\mathbf{z}_{(0)}, \mathbf{a}_{(0)}))_{r=1}^{kd_z}$ are linearly independent.

Notice that $\mathbf{z}_{(r)}$ represents values of $\mathbf{z}^{<t}$ for potentially different values of t and therefore can vary in shape.

The following example builds on Example 11 and shows that the sufficient variability of the above theorem might hold or not. The first case is interesting since it guarantees that $\mathbf{v}$ is linear while the second is interesting because it showcases a situation where the theory of Khemakhem et al. (2020a) and Lachapelle et al. (2022) do not apply (since they both rely on the above theorem), thus highlighting the importance of our nonparametric extension.

Example 15 (Satisfying or not the sufficient variability assumption of Theorem 4) We recall Example 11 in which $d_a = d_z$, $\mathbf{a} \in \mathcal{A} := \{\mathbf{0}, \mathbf{e}_1, \dots, \mathbf{e}_{d_a}\}$ and $\mathbf{G}^a = \mathbf{I}$ without temporal dependencies: For all $i \in [d_z]$, $p(\mathbf{z}_i | \mathbf{a}) = \mathcal{N}(\mathbf{z}_i; \boldsymbol{\mu}_i \mathbf{a}_i, 1 + \delta_i \mathbf{a}_i)$ where $\boldsymbol{\mu}_i \in \mathbb{R}$ and $\delta_i > -1$. We consider the cases where $\forall i, \delta_i = 0$ (variances do not change) and $\forall i, \delta_i \neq 0$ (variances change).

If $\forall i, \delta_i = 0$, we can represent $p(\mathbf{z}_i | \mathbf{a})$ in its exponential form with a one-dimensional sufficient statistic given by $\mathbf{s}_i(\mathbf{z}_i) = \mathbf{z}_i$ and a natural parameter given by $\boldsymbol{\lambda}_i(\mathbf{a}) = \boldsymbol{\mu}_i \mathbf{a}_i$. It can be easily seen that if $\forall i, \boldsymbol{\mu}_i \neq 0$ (i.e. the mean changes after the intervention), then the sufficient variability assumption of Theorem 4 holds since the vectors $\boldsymbol{\lambda}(\mathbf{e}_i) - \boldsymbol{\lambda}(\mathbf{0}) = \boldsymbol{\mu}_i \mathbf{e}_i$ span $\mathbb{R}^{d_z}$.

If $\forall i, \delta_i \neq 0$, we can represent $p(\mathbf{z}_i | \mathbf{a})$ in its exponential form with a two-dimensional sufficient statistics given by $\mathbf{s}_i(\mathbf{z}_i) = (\mathbf{z}_i, \mathbf{z}_i^2)$ and natural parameter given by $\boldsymbol{\lambda}_i(\mathbf{a}) = \left(\frac{\boldsymbol{\mu}_i \mathbf{a}_i}{1 + \delta_i \mathbf{a}_i}, \frac{-1}{2(1 + \delta_i \mathbf{a}_i)} \right)$. Note that, because we only have d_z interventions, for any choice of $\mathbf{a}_{(0)} \in \mathcal{A}$, the vectors $\{\boldsymbol{\lambda}(\mathbf{a}) - \boldsymbol{\lambda}(\mathbf{a}_{(0)})\}_{\mathbf{a} \in \mathcal{A}}$ can span at most a d_z -dimensional subspace, which is insufficient variability according to Theorem 4 since it requires spanning $\mathbb{R}^{2d_z}$.

4.3 Partial Disentanglement via Sparse Time Dependencies in Exponential Families

We now provide a (partial) disentanglement guarantee that takes advantage of sparsity regularization of $\hat{\mathbf{G}}^z$ and is specialized for exponential families with a one-dimensional sufficient statistic ($k = 1$). We will see that this extra parametric assumption on the transition model allows us to weaken the sufficient influence assumption of Theorem 3 (Assumption 8). In particular, this is going to allow for Markovian transitions without auxiliary variables, which was not allowed by the nonparametric result (Remark 7).

The sufficient influence assumption for $\mathbf{z}$ specialized to exponential families with $k = 1$ is directly taken from Lachapelle et al. (2022):

Assumption 11 (Sufficient influence of $\mathbf{z}$ (Lachapelle et al., 2022)) Assume $k = 1$ and $Ds(\mathbf{z})$ is invertible everywhere. There exist $\{(\mathbf{z}_{(r)}, \mathbf{a}_{(r)}, \tau_{(r)})\}_{r=1}^{\|\mathbf{G}^z\|_0}$ belonging to their respective support such that

$$\text{span} \left\{ D_z^{\tau_{(r)}} \boldsymbol{\lambda}(\mathbf{z}_{(r)}, \mathbf{a}_{(r)}) Ds(\mathbf{z}_{(r)}^{\tau_{(r)}})^{-1} \right\}_{r=1}^{\|\mathbf{G}^z\|_0} = \mathbb{R}_{\mathbf{G}^z}^{d_z \times d_z},$$

where $D_z^{\tau_{(r)}} \boldsymbol{\lambda}$ and Ds are Jacobians with respect to $\mathbf{z}^{\tau_{(r)}}$ and $\mathbf{z}$, respectively.

In Appendix B.4, we show that the above assumption is implied by its nonparametric version (Assumption 8) when the transition model is in an exponential family with $k = 1$. However, Assumption 11 is *strictly weaker* than its nonparametric counterpart, Assumption 8. The reason is that in the former, $\mathbf{z}^{\tau_{(r)}}$ can vary for different p , whereas this is not allowed in the latter, since we require $\mathbf{z} = \mathbf{z}^{\tau_{(r)}}$ for all r .

The following theorem, extended from [Lachapelle et al. \(2022\)](#), shows that making stronger parametric assumptions on the transition model allows to weaken the sufficient influence assumption. Note that its structure is nearly identical to Theorem 3. Its proof can be found in Appendix B.3.

Theorem 5 (Disentanglement via sparse temporal dependencies in exponential families) *Let $\theta := (\mathbf{f}, \lambda, \mathbf{G})$ and $\hat{\theta} := (\hat{\mathbf{f}}, \hat{\lambda}, \hat{\mathbf{G}})$ be two models satisfying Assumptions 1, 2, 3, 4, 9 as well as all assumptions of Theorem 4. Further suppose that*

1. *The sufficient statistic $\mathbf{s}$ is d_z -dimensional ($k = 1$) and is a diffeomorphism from $\mathbb{R}^{d_z}$ to $\mathbf{s}(\mathbb{R}^{d_z})$;*
2. *[Sufficient influence of $\mathbf{z}$] The Jacobian of the ground-truth transition function λ with respect to $\mathbf{z}$ varies “sufficiently”, as formalized in Assumption 11;*

Then, there exists a permutation matrix $\mathbf{P}$ such that $\mathbf{P}\mathbf{G}^z\mathbf{P}^\top \subseteq \hat{\mathbf{G}}^z$. Further assume that

3. *[Sparsity regularization] $\|\hat{\mathbf{G}}^z\|_0 \leq \|\mathbf{G}^z\|_0$;*

Then, $\theta \sim_{\text{con}}^z \hat{\theta}$ (Def. 14) & $\theta \sim_{\text{lin}} \hat{\theta}$ (Def. 17), which together implies that

$$\mathbf{v}(\mathbf{z}) = \mathbf{s}^{-1}(\mathbf{C}\mathbf{P}^\top \mathbf{s}(\mathbf{z}) + \mathbf{b}),$$

where $\mathbf{b} \in \mathbb{R}^{d_z}$ and $\mathbf{C} \in \mathbb{R}^{d_z \times d_z}$ is invertible, $\mathbf{G}^z$ - and $(\mathbf{G}^z)^\top$ -preserving (Definition 11).

The reason we can simplify the sufficient influence assumption in the exponential family case has to do with the quasi-linear form of $\mathbf{v}$. Indeed, in that case, one can compute that the Jacobian of $\mathbf{v}$ takes a special form: $D\mathbf{v}(\mathbf{z}) = D\mathbf{s}(\mathbf{v}(\mathbf{z}))^{-1}\mathbf{L}D\mathbf{s}(\mathbf{z})$. Since $D\mathbf{s}$ is diagonal everywhere, one can see that the “non-diagonal part” of $D\mathbf{v}(\mathbf{z})$, i.e. $\mathbf{L}$, does not depend on $\mathbf{z}$, which simplifies the proof. See Appendix B.3 for more details.

In Appendix D.2, we discuss how [Khemakhem et al. \(2020a\)](#) & [Yao et al. \(2022b\)](#) obtain disentanglement guarantees and how their assumptions differ from ours.

Example 16 (Markovian sparse temporal dependencies without auxiliary variables) *Recall Remark 7 which pointed out that, without auxiliary variables, non-Markovianity was necessary to satisfy the nonparametric Assumption 8. We now illustrate that the analogous assumption specialized for exponential families with $k = 1$ (Assumption 11) is not as restrictive, i.e. it allows for Markovianity even when there are no auxiliary variables.*

We start from Example 7 which was based on the situation depicted in Figures 1 & 3e where the temporal graph $\mathbf{G}^z$ is lower triangular. Assume that no action is observed, i.e. we can only leverage the sparsity of $\mathbf{G}^z$ to disentangle. We now provide a concrete Markovian transition model $p(\mathbf{z}^t \mid \mathbf{z}^{t-1})$ that satisfies Assumption 11. Similarly to previous examples, assume $p(\mathbf{z}^t \mid \mathbf{z}^{t-1}) = \mathcal{N}(\mathbf{z}^t; \boldsymbol{\mu}(\mathbf{z}^{t-1}), \sigma^2 \mathbf{I})$ where

$$\boldsymbol{\mu}(\mathbf{z}) := \mathbf{z} + \begin{bmatrix} z_1^2/2 \\ z_1^3/3 \\ z_1^4/4 \end{bmatrix} + \begin{bmatrix} 0 \\ z_2^2/2 \\ z_2^3/3 \end{bmatrix} + \begin{bmatrix} 0 \\ 0 \\ z_3^2/2 \end{bmatrix}.$$

Because the variance σ^2 is not influenced by $\mathbf{z}^{t-1}$, we can represent this transition model in an exponential family with $k = 1$ where the natural parameter is given by

$$\boldsymbol{\lambda}(\mathbf{z}^{t-1}) = \mu(\mathbf{z}^{t-1})/\sigma$$

and the sufficient statistic is given by $\mathbf{s}(\mathbf{z}) = \mathbf{z}/\sigma$. We can thus compute

$$D\boldsymbol{\lambda}(\mathbf{z})D\mathbf{s}(\mathbf{z})^{-1} = \mathbf{I} + \begin{bmatrix} z_1 & 0 & 0 \\ z_1^2 & z_2 & 0 \\ z_1^3 & z_2^2 & z_3 \end{bmatrix}$$

which spans the 6-dimensional space $\mathbb{R}_{\mathbf{G}^z}^{3 \times 3}$, as shown in Example 14.

Connecting with the sufficient influence assumption of $\mathbf{a}$ in Lachapelle et al. (2022) The previous work of Lachapelle et al. (2022) could also leverage the sparse influence of $\mathbf{a}$ to disentangle and was based on exponential family and sufficient influence assumptions. In Appendix B.4, Proposition 12 shows that their sufficient influence assumption for $\mathbf{a}$ is actually equivalent to our nonparametric version (Assumption 7) in the exponential family case with $k = 1$. An important conclusion of this observation is that the identifiability result via sparse $\mathbf{G}^a$ from Lachapelle et al. (2022), which was limited to the exponential family case with $k = 1$, can be derived from the more general nonparametric result of Theorem 2 we introduced earlier.

5 Model Estimation with a Sparsity Constraint

The identifiability results presented in this work are based on two crucial postulates: (i) the distribution over observations of both the learned and ground-truth models must match, i.e. $\hat{\boldsymbol{\theta}} \sim_{\text{obs}} \boldsymbol{\theta}$ (Definition 4), and (ii) the learned graphs $\hat{\mathbf{G}}^a$ and $\hat{\mathbf{G}}^z$ must be as sparse as their ground-truth counterparts, respectively $\mathbf{G}^a$ and $\mathbf{G}^z$. The theory suggests that in order to learn a (partially) disentangled representation, one should learn a model that satisfies these two requirements. In this section, we present one particular practical approach to achieve this. Appendix C.2 provides further details.

Data fitting. The first condition can be achieved by fitting a model to the data. Since the models discussed in this work present latent variable models, it is a natural idea to use a maximum likelihood approach based on the well-known framework of variational autoencoders (VAEs) (Kingma and Welling, 2014) in which the decoder neural network corresponds to the mixing function $\hat{\mathbf{f}}$. We consider an approximate posterior of the form

$$q(\mathbf{z}^{\leq T} \mid \mathbf{x}^{\leq T}, \mathbf{a}^{\leq T}) := \prod_{t=1}^T q(\mathbf{z}^t \mid \mathbf{x}^t), \quad (32)$$

where $q(\mathbf{z}^t \mid \mathbf{x}^t)$ is a Gaussian distribution with mean and diagonal covariance outputted by a neural network $\text{encoder}(\mathbf{x}^t)$. In our experiments, the latent model $\hat{p}(\mathbf{z}_i^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t})$ is a Gaussian distribution with mean $\hat{\mu}_i(\mathbf{z}^{<t}, \mathbf{a}^{<t})$ parameterized as a fully connected neural network that “looks” only at a fixed window of s lagged latent variables.⁷ Furthermore, the variances are learned but does

7. The theory we developed would allow for a μ function that depends on all previous time steps, not only the s previous ones. This could be achieved with a recurrent neural network or transformer, but we leave this to future work.

not depend on $(\mathbf{z}^{<t}, \mathbf{a}^{<t})$ (see Appendix C.2 for details). This variational inference model induces the following evidence lower bound (ELBO) on $\log \hat{p}(\mathbf{x}^{\leq T} | \mathbf{a}^{<T})$:

$$\log \hat{p}(\mathbf{x}^{\leq T} | \mathbf{a}^{<T}) \geq \text{ELBO}(\hat{\mathbf{f}}, \hat{\boldsymbol{\mu}}, \hat{\mathbf{G}}, q; \mathbf{x}^{\leq T}, \mathbf{a}^{<T}) := \sum_{t=1}^T \mathbb{E}_{q(\mathbf{z}^t | \mathbf{x}^t)} [\log \hat{p}(\mathbf{x}^t | \mathbf{z}^t)] - \mathbb{E}_{q(\mathbf{z}^{<t} | \mathbf{x}^{<t})} KL(q(\mathbf{z}^t | \mathbf{x}^t) || \hat{p}(\mathbf{z}^t | \mathbf{z}^{<t}, \mathbf{a}^{<t})). \quad (33)$$

We derive this fact in Appendix D.3. This lower bound can then be maximized using some variant of stochastic gradient ascent such as Adam (Kingma and Ba, 2015). We note that many works have proposed learning dynamical models with latent variables using VAEs (Girin et al., 2020), with various choice of architectures and approximate posteriors. Our specific choices were made out of a desire for simplicity, but the reader should be aware of other possibilities.

The learned distribution will exactly match the ground truth distribution if (i) the model has enough capacity to express the ground-truth generative process, (ii) the approximate posterior has enough capacity to express the ground-truth posterior $p(\mathbf{z}^t | \mathbf{x}^{\leq T}, \mathbf{a}^{<T})$, (iii) the dataset is sufficiently large, and (iv) the optimization finds the global optimum. If, in addition, the ground truth generative process satisfies the assumptions of Proposition 2, we can guarantee that the learned model $\hat{\boldsymbol{\theta}}$ will be equivalent to the ground truth model $\boldsymbol{\theta}$ up to diffeomorphism (Definition 5).

Learning $\hat{\mathbf{G}}$ with sparsity constraints. To go from equivalence up to diffeomorphism to actual disentanglement (partial or not), Theorems 1, 2, 3 & 5 suggest that we should not only fit the data, but also choose the learned graph $\hat{\mathbf{G}}$ such that $\|\hat{\mathbf{G}}^a\|_0 \leq \|\mathbf{G}^a\|_0$ and/or $\|\hat{\mathbf{G}}^z\|_0 \leq \|\mathbf{G}^z\|_0$. In order to allow gradient-based optimization, our strategy consists in treating each edge $\hat{\mathbf{G}}_{i,j}$ as an independent Bernoulli random variable with probability of success $\sigma(\gamma_{i,j})$, where σ is the sigmoid function and $\gamma_{i,j}$ is a parameter learned using the Gumbel-Softmax trick (Jang et al., 2017; Maddison et al., 2017). Let $\text{ELBO}(\hat{\mathbf{f}}, \hat{\boldsymbol{\mu}}, \hat{\mathbf{G}}, q)$ be the ELBO objective of (33) averaged over the whole dataset. We tackle the following constrained optimization problem:

$$\max_{\hat{\mathbf{f}}, \hat{\boldsymbol{\mu}}, \gamma, q} \mathbb{E}_{\hat{\mathbf{G}} \sim \sigma(\gamma)} \text{ELBO}(\hat{\mathbf{f}}, \hat{\boldsymbol{\mu}}, \hat{\mathbf{G}}, q) \quad \text{subject to} \quad \mathbb{E}_{\hat{\mathbf{G}} \sim \sigma(\gamma)} \|\hat{\mathbf{G}}\|_0 \leq \beta.$$

where β is a hyperparameter (which should ideally be set to $\beta^* := \|\mathbf{G}\|_0$, i.e., the number of edges in the ground-truth graph) and $\hat{\mathbf{G}} \sim \sigma(\gamma)$ means that $\hat{\mathbf{G}}_{i,j}$ are independent and distributed according to $\sigma(\gamma_{i,j})$. Because $\mathbb{E}_{\hat{\mathbf{G}} \sim \sigma(\gamma)} \|\hat{\mathbf{G}}\|_0 = \|\sigma(\gamma)\|_1$ where $\sigma(\gamma)$ is a matrix, the constraint becomes $\|\sigma(\gamma)\|_1 \leq \beta$. To solve this problem, we perform gradient descent-ascent on the Lagrangian function given by

$$\mathbb{E}_{\hat{\mathbf{G}} \sim \sigma(\gamma)} \text{ELBO}(\hat{\mathbf{f}}, \hat{\boldsymbol{\mu}}, \hat{\mathbf{G}}, q) - \alpha(\|\sigma(\gamma)\|_1 - \beta)$$

where the ascent step is performed w.r.t. $\hat{\mathbf{f}}, \hat{\boldsymbol{\mu}}, \hat{\mathbf{G}}$ and q ; and the descent step is performed w.r.t. Lagrangian multiplier α , which is forced to remain greater or equal to zero via a simple projection step. As suggested by Gallego-Posada et al. (2021), we perform *dual restarts* which simply means that, as soon as the constraint is satisfied, the Lagrangian multiplier is reset to 0. We used the `Cooper` library (Gallego-Posada and Ramirez, 2022), which implements many constrained optimization procedures in Python, including the one described above. Note that we use Adam (Kingma and Ba, 2015) for the ascent steps and standard gradient descent for the descent step on the Lagrangian multiplier α .

We also found empirically that the following schedule for β is helpful: We start training with $\beta = \max_{\hat{G}} \|\hat{G}\|_0$ and linearly decrease its value until the desired number of edges is reached. This avoids getting a sparse graph too quickly while training, thus giving enough time to the model parameters to adjust. In each experiment, we trained for 300K iterations, and β takes 150K to go from its initial value to its desired value. We discuss how to select the hyperparameter β in Section 8.

6 Evaluation with R_{con} and SHD

In this section, we tackle the problem of quantitatively evaluating whether a learned representation $\hat{z}$ is completely or partially disentangled with respect to the ground-truth representation z , given a dataset of paired representations $\{(z^i, \hat{z}^i)\}_{i \in [N]}$. More precisely, we want to evaluate whether two models are α -consistent or z -consistent (Definitions 13 & 14). To achieve this, we must evaluate whether there exists a graph preserving map c (Definition 12) and a permutation matrix P such that for all $i \in [N]$, $z^i = c(P^\top \hat{z}^i)$. For evaluation purposes, we assume we observe the ground-truth latent representation for each observation, i.e., we have $\{(x^i, z^i)\}_{i \in N}$ sampled i.i.d. from the ground-truth data generating process. We will take $\hat{z}^i := \text{encoder}(x^i)$ where encoder is from the learned VAE model introduced in Section 5. For simplicity, we assume that c is affine.⁸

We start with evaluating complete disentanglement. A popular choice for this is the *mean correlation coefficient* (MCC), which is obtained by first computing the Pearson correlation matrix $K \in \mathbb{R}^{d_z \times d_z}$ between the ground-truth representation and the learned representation ($K_{i,j}$ is the correlation between z_i and $\hat{z}_j$). Then $\text{MCC} := \max_{P \in \text{permutations}} \frac{1}{d_z} \sum_{i=1}^{d_z} |(KP)_{i,i}|$. We denote by $\hat{P}$ the optimal permutation found by MCC.

To evaluate whether the learned representation is identified up to linear transformation (Definition 17), we perform linear regression to predict the ground-truth latent factors from the learned ones, and report the mean of the Pearson correlations between the predicted ground-truth latents and the actual ones. This metric is sometimes called the *coefficient of multiple correlation*, and happens to be the square root of the better known *coefficient of determination*, usually denoted by R^2 . The advantage of using R instead of R^2 is that the former is comparable to MCC, and we always have $\text{MCC} \leq R$. Let us denote by $\hat{L}$ the matrix of estimated coefficients, which should be thought of as an estimation of L in Definition 17 (assuming $s(z) = z$, as is the case with Gaussian latents with fixed variance). Note that $\hat{L}$ was fitted to the standardized z and $\hat{z}$ (shifted and scaled to have mean 0 and variance 1). This yields coefficients $\hat{L}_{i,j}$ that are directly comparable without changing the value of the R score. We visualize $\hat{L}$ in Figures 6 & 8.

To evaluate whether the learned representation is α -consistent or z -consistent with the ground-truth (Definitions 13 & 14), as predicted by Theorems 1 & 3, we introduce a new metric, denoted by R_{con} . The idea behind R_{con} is to predict the ground-truth factors z from only the inferred factors $\hat{z}$ that are allowed by the equivalence relations. For example, for α -consistency (Definition 13), the relation between z and $\hat{z}$ is given by $z = c(P^\top \hat{z})$ where c is a G^a -preserving diffeomorphism. Since we assume for simplicity that c is affine, we have $z = CP^\top \hat{z} + b$ where C is a G^a -preserving matrix. The idea is then to estimate both P and C using samples $(z, \hat{z})$. The permutation P is estimated by $\hat{P}$, which was found when computing MCC (Section 8). To estimate C , we compute $\hat{z}_{\text{perm}} := \hat{P}^\top \hat{z}$ and then compute the mask $M \in \{0, 1\}^{d_z \times d_z}$ specifying which entries of C are

8. This is not a simplification when the latent factors in the model and in the data-generating process are Gaussian with fixed variance and the assumptions of Theorem 4 hold. That is because the latent model is in the exponential family with sufficient statistic $s(z) = z$ and, by Theorem 4, we must have $z = s^{-1}(Ls(\hat{z}) + b) = L\hat{z} + b$.

allowed to be nonzero, as required by the G^a -preservation property (Proposition 3). Then, for every i , we predict the ground-truth z_i by performing linear regression only on the *allowed* factors, i.e., $M_{i\cdot} \odot \hat{z}_{\text{perm}}$, and compute the associated coefficient of multiple correlations $R_{\text{con},i}$ and report the mean, i.e., $R_{\text{con}} := \frac{1}{d_z} \sum_{i=1}^{d_z} R_{\text{con},i}$. It is easy to see that we must have $R_{\text{con}} \leq R$, since R_{con} was computed with fewer features than R . Moreover, $\text{MCC} \leq R_{\text{con}}$, because MCC can be thought of as computing exactly the same thing as for R_{con} , but by predicting z_i only from $\hat{z}_{\text{perm},i}$, i.e. with fewer features than R_{con} . This means that we always have $0 \leq \text{MCC} \leq R_{\text{con}} \leq R \leq 1$. This is a nice property allowing us to compare all three metrics together and reflects the hierarchy between equivalence relations. Note that R_{con} depends implicitly on the ground-truth graph, since the matrix M indicating which entries of C are forced to be zero by the equivalence relation depends on G .

To compare the learned graph $\hat{G}$ with the ground-truth G , we report the (normalized) *structural Hamming distance* (SHD) between the ground-truth graph and the estimated graph *permuted* by $\hat{P}$. More precisely, we report $\text{SHD} = (\|G^a - \hat{P}^\top \hat{G}^a\|_0 + \|G^z - \hat{P}^\top \hat{G}^z \hat{P}\|_0) / (d_a d_z + d_z^2)$, where $\hat{P}$ is the permutation found by MCC and $(d_a d_z + d_z^2)$ is the maximal number of edges G can have.

7 Related Work

Causal discovery. Causal discovery is the problem of learning a causal graph from observational and possibly interventional data when all variables are observed, unlike in the present work where the representation z is latent. In this simpler setting, graph identifiability is still challenging since, without interventions, one can only identify the Markov equivalence class of the causal graph (thus leaving some edge orientation ambiguous), provided the distribution entailed by the causal graphical model is faithful to its causal graph, i.e. its conditional independences are precisely the ones induced by the causal graph. The faithfulness condition implies that we can infer precisely which d -separation statements hold in the causal graph by performing conditional independence tests on the observed distribution. This class of methods is known as *constrained-based* (Pearl, 2009; Peters et al., 2017). Another approach is to define a score function on the space of graphs, usually a sparsity-regularized maximum likelihood score, and search the space of graphs to maximize it. These are known as *score-based* methods (Chickering, 2003; Shimizu et al., 2006; Peters et al., 2014; Zheng et al., 2018). These approaches can be specialized for the interventional setting (Hauser and Bühlmann, 2012; Brouillard et al., 2020) and the temporal setting (Pamfil et al., 2020; Runge et al., 2019).

Linear and nonlinear ICA. The first results showing that latent variables can be identified up to permutation and rescaling at least date back to classical linear ICA which assumes a linear mixing function f and mutually independent and non-Gaussian latent variables (Jutten and Herault, 1991; Tong et al., 1993; Comon, 1994). Hyvärinen and Pajunen (1999) showed that when allowing f to be a general nonlinear transformation, a setting known as nonlinear ICA, mutual independence and non-Gaussianity alone are insufficient to identify the latent variables. This inspired multiple variations of nonlinear ICA that enabled identifiability by leveraging, e.g., nonstationarity (Hyvärinen and Morioka, 2016) and temporal dependencies (Hyvärinen and Morioka, 2017). Hyvärinen et al. (2019) generalized these works by introducing a data generating process in which latent variables are conditionally mutually independent given an observed *auxiliary variable* (corresponding to a in our work). These last three works rely on some form of *noise contrastive estimation* (NCE) (Gutmann and Hyvärinen, 2012), but similar identifiability results have also been shown for VAEs (Khe-

makhem et al., 2020a; Locatello et al., 2020; Klindt et al., 2021), normalizing flows (Sorrenson et al., 2020) and energy-based models (Khemakhem et al., 2020b). Kivva et al. (2022) showed that it is not necessary to observe the auxiliary variable to obtain disentanglement when the mixing function is piecewise affine and the latent factors are distributed according to a mixture of Gaussians with diagonal covariance.

Causal representation learning (static). Since the publication of the first iteration of this work in CLear 2022, the field now known as *causal representation learning* (CRL) (Schölkopf et al., 2021) gained significant traction. The prototypical problem of CRL is similar to nonlinear ICA in that the goal is to identify latent factors of variations but differs in that the latent variables are assumed to be related via a causal graphical model (CGM) and interventions on the latents are typically observed. Although some works assumed that the structure of the causal graph is known (Kocaoglu et al., 2018; Shen et al., 2022; Nair et al., 2019; Liang et al., 2023), significant progress has been achieved recently in the setting where the latent causal graph is unknown and must be inferred from single-node interventions targeting latent variables (Ahuja et al., 2023; Squires et al., 2023; Buchholz et al., 2023; von Kügelgen et al., 2023; Zhang et al., 2023; Jiang and Aragam, 2023; Varici et al., 2023b,a). See Remark 4 for a discussion contrasting these works with our interventional framework. In a similar spirit, Liu et al. (2023); Yang et al. (2021) leverage a form of nonstationarity that does not necessarily correspond to interventions, and Bengio et al. (2020) suggests using adaptation speed as a heuristic objective to disentangle latent factors in the bivariate case, although without identifiability guarantees. The above works do not support temporal dependencies, unlike the framework presented in this work. Although we do focus on temporal dependencies, the special case where $T = 1$ discussed in Remark 4 and fleshed out in Examples 10, 11 & 13 can be categorized as static CRL where the latent factors are independent (Assumption 2). This approach has been applied to single-cell data with gene perturbations (Lopez et al., 2023; Bereket and Karaletsos, 2023). Importantly, Example 13 illustrates how *multi-node* interventions on the latent factors can yield (partial) disentanglement in the independent factors regime. To the best of our knowledge, this constitutes the first identifiability guarantee from multi-node interventions with nonlinear mixing and should form an important step towards generalizing to arbitrary latent graphs. Note that Bing et al. (2023) recently proposed a disentanglement guarantee from multi-node interventions in the linear mixing setting.

CRL is closely related to methods that assume access to **paired observations** (x, x') that are generated from a common decoder f . These are in contrast with the works discussed above, which assume the samples from observational and interventional distributions are **unpaired**. In the *paired* data regime, Locatello et al. (2020) and Ahuja et al. (2022b) assume that only a small set of latent factors $S \subseteq [d_z]$ changes between x and x' . Interestingly, Locatello et al. (2020) assume that for all i , $P(S \cap S' = \{i\}) > 0$ (for i.i.d S and S'), which resembles our graphical criterion for complete disentanglement (Definition 5). Karaletsos et al. (2016) proposed a related strategy based on triplets of observations and weak labels indicating which observation is closer to the reference in the (masked) latent space. Von Kügelgen et al. (2021) modeled the self-supervised setting with data augmentation using a similar idea and showed the block-identifiability of the latent variables shared among x and x' . Brehmer et al. (2022) assumes that the latent variables are sampled from a structural causal model (SCM) (Peters et al., 2017) and that x' is *counterfactual* in the sense that it is generated using the same SCM and exogenous noise values as x except for some noises which are modified randomly. Similar approaches can also provide identifiability guarantees in the *multi-*

view setting where the decoders for the different views $\mathbf{x}$ and $\mathbf{x}'$ are allowed to be different (Gresele et al., 2020; Daunhawer et al., 2023). The paired observation setting bears some similarity with the temporal setting covered in this work since the pairs $(\mathbf{x}^t, \mathbf{x}^{t-1})$ are observed jointly. However, contrary to the above works, Theorems 3 & 5 allow all latent variables to change between $t - 1$ and t , only the temporal dependencies between them are assumed to be sparse. Morioka and Hyvarinen (2023) can also be seen as paired CRL in which the latents of different views can interact causally in a restricted manner. Recently, Yao et al. (2023) generalized previous work by allowing more than two views.

Leveraging temporal dependencies or non-stationarity. Tong et al. (1990) proved the identifiability of linear ICA when latent factors z_i^t are correlated across time steps t but remain independent across components i , an idea that has been extended to nonlinear mixing (Hyvarinen and Morioka, 2017; Klindt et al., 2021; Schell and Oberhauser, 2023). Using our notation, these works assume a diagonal adjacency matrix $\mathbf{G}^z$ which contrasts with Theorems 3 & 5 which allow for general $\mathbf{G}^z$ (although some graphs might not yield complete disentanglement). Yao et al. (2022a, Theorem 1) also allows for general $\mathbf{G}^z$, but does not rely on the sparsity of $\mathbf{G}^z$ or sparse interventions on latent factors for identification. Instead, it relies on the conditional independence of z_i^t given z^{t-1} and on a “sufficient variability” condition involving the third cross-derivatives $\frac{\partial^3}{(\partial z_i^t)^2 \partial z_j^{t-1}} \log p(z_i^t | z^{t-1})$ which excludes simple Gaussian models with homoscedastic variance like the ones we considered in Examples 9, 10, 12 and in our experiments of Section 8. General non-stationarity of the latent distribution, i.e., that are not necessarily sparse like what is considered in this work, can also be used to identify the latent factors (Hyvarinen and Morioka, 2016; Hyvärinen et al., 2019; Khemakhem et al., 2020a; Hälvä and Hyvärinen, 2020; Morioka et al., 2021; Yao et al., 2022b,a), but these results require sufficient variability of higher-order derivatives/differences of the log-densities, which again typically exclude simple homoscedastic Gaussian models (see Appendix D.2 for more). Ahuja et al. (2022a) characterized the indeterminacies of the representation in latent dynamical models as the set of equivariances of the transition mechanism. In addition to temporal dependencies, one can also consider latent factors structured according to a spatial topology (Hälvä et al., 2021).

Dynamical causal representation learning: The previous iteration of this work (Lachapelle et al., 2022) concurrently with Lippe et al. (2022) introduced latent variables identifiability guarantees for dynamical latent models based on sparse interventions. Lippe et al. (2023b) later proposed a generalization in which instantaneous causal connections are allowed. The key differences from the present work are (i) Lippe et al. (2023b) considers interventions with *known targets*, while the present work (as well as Lachapelle et al. (2022)) considers interventions with *unknown targets*; (ii) Lachapelle et al. (2022, Theorem 5) and Theorems 3 & 5 do not necessarily need interventions to disentangle since they leverage sparsity of the temporal dependencies, contrary to Lippe et al. (2023b); (iii) Lippe et al. (2023b) allows instantaneous causal connections, unlike the present work; and (iv) Lippe et al. (2023b) demonstrates their approach on image data. We hypothesize that the *partially-perfect interventions* from Lippe et al. (2023b) in which instantaneous causal effects are disabled might be enough to get identifiability in our framework. We leave this for future work. The concurrent work of Volodin (2021) independently proposed an approach very similar to ours, which also learns a sparse latent causal graph over latent variables and actions using binary masks, but focuses on testing various algorithmic variants and empirically verifies that the approach works on interactive environments rather than formal identifiability guarantees. Lopez et al. (2023); Lei et al. (2023) found that such models adapt to sparse interventions more quickly than their entangled

counterparts. [Keurti et al. \(2023\)](#) discusses disentanglement in temporal regimes through the lens of group theory but does not provide identifiability guarantees. Recently, [Lippe et al. \(2023a\)](#) proposed a model similar to ours with disentanglement guarantees based on the constraint that the effect of the variable α^{t-1} (analogous to R in their work) on each z_i^t is mediated by a deterministic binary variable.

Constraining the decoder function f . It is worth noting that one can also obtain disentanglement guarantees by constraining the decoder function f in some way ([Taleb and Jutten, 1999](#); [Gresele et al., 2021](#); [Buchholz et al., 2022](#); [Leemann et al., 2023](#); [Lachapelle et al., 2023b](#); [Horan et al., 2021](#)). In particular, this can be achieved by imposing some form of sparsity on f ([Moran et al., 2022](#); [Zheng et al., 2022](#); [Brady et al., 2023](#); [Xi and Bloem-Reddy, 2023](#)). In contrast, the present work assumes only that f is a general diffeomorphism onto its image. Note that [Zheng et al. \(2022\)](#) reused many proof strategies from the shorter version of this work ([Lachapelle et al., 2022](#)).

Disentanglement with explicit supervision. Some works leverage more explicit supervision to disentangle. For example, [Ahuja et al. \(2022c\)](#) assumes that the labels are given by a linear transformation of mutually independent and non-Gaussian latent factors. Instead of relying on independence, [Lachapelle et al. \(2023a\)](#); [Fumero et al. \(2023\)](#) leverage the sparsity of the linear map to disentangle.

Other relevant works on sparsity. The assumption that high-level variables are sparsely related to each other and/or to actions was discussed by [Bengio \(2019\)](#); [Goyal and Bengio \(2021\)](#); [Ke et al. \(2021\)](#). These ideas have also been leveraged by [Goyal et al. \(2021b,a\)](#); [Madan et al. \(2021\)](#) via attention mechanisms. Although these works are, in part, motivated by the same core assumption as ours, their focus is more on empirically verifying out-of-distribution generalization than it is on disentanglement (Definition 7) and formal identifiability results. The assumption that individual actions often affect only one factor of variation has been leveraged to disentangle [Thomas et al. \(2017\)](#). Loosely speaking, the theory we developed in the present work can be seen as a formal justification for such approaches.

8 Experiments

To illustrate our identifiability results and the benefit of mechanism sparsity regularization for disentanglement, we apply the sparsity regularized VAE method of Section 5 on various synthetic datasets. Section 8.1 focuses on graphs satisfying the criterion of Assumption 5 which, as we saw, guarantees complete disentanglement. We also verify experimentally that the sufficient influence assumptions are indeed important for disentanglement and explore latent model with both homoscedastic and heteroscedastic variance. Section 8.2 explores graphs that do not satisfy the criterion. Details about our implementation are provided in Appendix C.2 and the code used to run these experiments can be found here: https://github.com/slachapelle/disentanglement_via_mechanism_sparsity.

Synthetic datasets. The datasets we considered are separated into two groups: *Action & Time* datasets. The former group has only auxiliary variables, which we interpret as actions, without temporal dependence, so we fix $\hat{G}^z = \mathbf{0}$. The latter group has only temporal dependence without actions, so we fix $\hat{G}^a = \mathbf{0}$. In each dataset, the ground-truth mixing function f is a randomly initialized neural network. The dimensionalities of z and x are $d_z = 10$ and $d_x = 20$, respec-

tively. In the action datasets, the dimensionality of $\mathbf{a}$ is $d_a = 10$, unless otherwise specified. The ground-truth transition model $p(\mathbf{z}^t | \mathbf{z}^{<t}, \mathbf{a}^{<t})$ is always a Gaussian with a mean outputted by some function $\mu_G(\mathbf{z}^{t-1}, \mathbf{a}^{t-1})$ (the data is *Markovian*). For all datasets considered, the covariance matrix is given by $\sigma_z^2 \mathbf{I}$, i.e., the variance is *homoscedastic*, except for the datasets ActionNonDiag $_{k=2}$ and TimeNonDiag $_{k=2}$, which have *heteroscedastic* variance. Appendix C.1 provides more detailed descriptions of the datasets including the explicit form of μ and G in each case. Note that the learned transition model $\hat{p}(\mathbf{z}^t | \mathbf{z}^{t-1}, \mathbf{a}^{t-1})$ is also a homoscedastic Gaussian where the mean function $\hat{\mu}$ is an MLP.

Baselines. On the action datasets, we compare with TCVAE (Chen et al., 2018), iVAE (Khemakhem et al., 2020a). Only iVAE leverages the action. On the temporal datasets, we compare our approach with TCVAE, PCL (Hyvarinen and Morioka, 2017) and SlowVAE (Klindt et al., 2021). Only PCL and SlowVAE leverage the temporal dependencies. We also report the performance of a randomly initialized encoder (Random) and one trained via least-square regression directly on the ground-truth latent factors (Supervised). See Appendix C.3 for details on the baselines.

Unsupervised hyperparameter selection. In practice, the hyperparameters cannot be selected so as to optimize MCC, since this metric requires access to the ground-truth latent factors. Duan et al. (2020) introduced *unsupervised disentanglement ranking* (UDR) as a solution to unsupervised hyperparameter selection for disentanglement. Figures 5 & 7 show the performance of all approaches using UDR to select the hyperparameter (when it has one). For our approach, we show a range of sparsity bounds β and indicate the hyperparameter selected by UDR with a black star. Note that for our approach, we excluded from the UDR selection hyperparameters that yielded graphs with fewer edges than latent factors as a heuristic to prevent UDR from selecting overly sparse graphs. Figures 5 & 7 show that this *unsupervised* procedure selects a reasonable regularization level (as indicated by the black star), although not always the optimal one. See Appendix C.4 for more details.

8.1 Graphs Allowing Complete Disentanglement (Satisfying Assumption 5)

Satisfying sufficient influence assumptions. Figure 5 reports the MCC and R scores of all methods on four datasets that satisfy both the graphical criterion and the sufficient influence assumption: the datasets ActionDiag and TimeDiag have “diagonal” graphs, i.e., $G^a = \mathbf{I}$ and $G^z = \mathbf{I}$, while ActionNonDiag and TimeNonDiag present more involved graphs (depicted in Figure 6).

Observations: We see that the sparsity constraint improves MCC and SHD on all datasets. Although most baselines obtain good R scores, which indicates that their representation encodes all the information about the factors of variations, they obtain poor MCC in comparison to our approach with a properly selected sparsity level, which indicates that they fail to disentangle. Moreover, the sparsity level selected by UDR (indicated by a black star) corresponds to the lowest SHD value for three out of four datasets and when it does not, it is still better than no sparsity at all. Figure 6 shows examples of estimated graphs. More details can be found in the caption.

Violating sufficient influence assumptions. The left column of Table 3 reports the performance of all methods on the ActionNonDiag $_{\text{NoSuffInf}}$ and TimeNonDiag $_{\text{NoSuffInf}}$ datasets, which are essentially the same as ActionNonDiag and TimeNonDiag but do not satisfy the sufficient influence assumptions (see Appendix C.1 for details). **Observations:** For the ActionNonDiag $_{\text{NoSuffInf}}$ dataset, we still see an improvement in MCC and SHD when regularizing for sparsity, but not as important

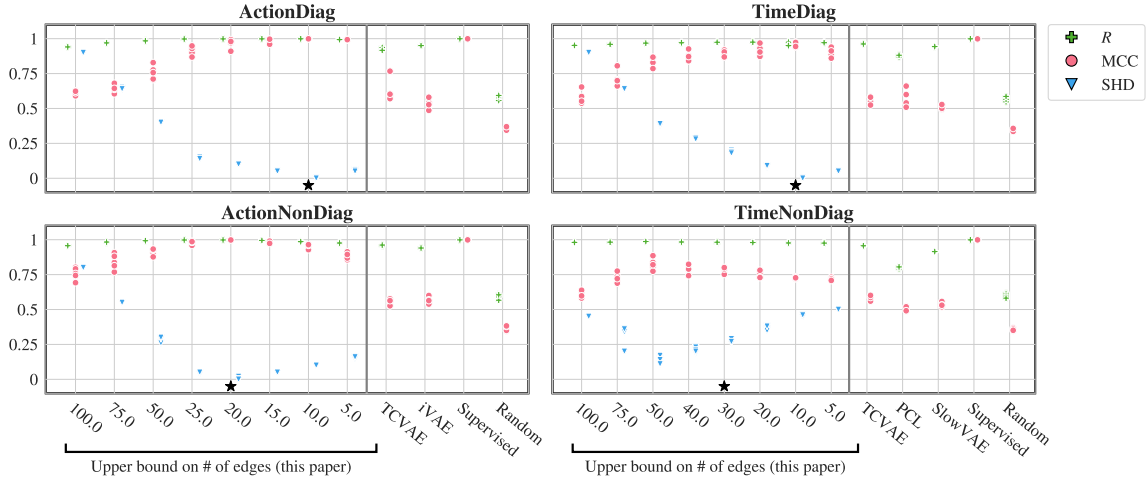


Figure 5: **Graphical criterion holds:** Datasets ActionDiag and TimeDiag have diagonal graphs while ActionNonDiag and TimeNonDiag have non-diagonal graphs. Sufficient influence is always satisfied. For our regularized VAE approach, we report performance for multiple sparsity levels β . In the left column, only $\hat{G}^a$ is learned while in the right column, only $\hat{G}^z$ is learned. For more details on the synthetic datasets, see Appendix C.1. The black star indicates which regularization parameter is selected by the filtered UDR procedure (see Appendix C.4). For R and MCC , higher is better. For SHD , lower is better. Performance is reported on 5 random seeds.

as for ActionNonDiag, which got $MCC \approx 1$ and $SHD \approx 0$. Nevertheless, our approach outperforms the baselines. For the TimeNonDiag_{NoSuffInf} dataset, there is simply no improvement in MCC from sparsity regularization. In that case, SlowVAE (with the hyperparameter selected to maximize MCC) and PCL have a higher MCC . These observations confirm the importance of the sufficient influence assumptions.

Heteroscedastic variance ($k = 2$). The right column of Table 3 reports the performance of all methods on the ActionNonDiag _{$k=2$} and TimeNonDiag _{$k=2$} datasets, which are essentially the same as ActionNonDiag and TimeNonDiag but presents heteroscedastic variance, i.e., $\text{var}(\mathbf{z}^t \mid \mathbf{z}^{t-1}, \mathbf{a}^{t-1})$ is not a constant function of $(\mathbf{z}^{t-1}, \mathbf{a}^{t-1})$. This setting is interesting since it is not covered by the exponential family theory of Lachapelle et al. (2022) which assumed a one-dimensional sufficient statistic $\mathbf{s}$ ($k = 1$) whereas here we have $k = 2$. Both datasets fall under the umbrella of our nonparametric theory. However, TimeNonDiag _{$k=2$} cannot satisfy the sufficient influence assumption because the data is Markovian and do not present an auxiliary variable (see Remark 7). **Observations:** Both datasets benefit from sparsity and outperform baselines. On ActionNonDiag _{$k=2$} we obtain near perfect MCC and SHD while on TimeNonDiag _{$k=2$} we obtain a performance similar to TimeNonDiag. We hypothesize that the performance bottleneck in both TimeNonDiag and TimeNonDiag _{$k=2$} is the estimation of the graph, as in both cases SHD is always greater than $\approx 20\%$.

Datasets	ActionNonDiag _{NoSuffInf}			ActionNonDiag _{k=2}		
Metrics	SHD	MCC	R	SHD	MCC	R
iVAE	–	.61±.02	.97±.00	–	.59±.03	.94±.00
TCVAE (UDR)	–	.58±.03	.88±.01	–	.55±.02	.96±.00
TCVAE (MCC)	–	.61±.02	.96±.00	–	.55±.02	.96±.00
Ours (no sparsity)	.80±.00	.62±.02	.93±.00	.80±.00	.70±.03	.97±.00
Ours (sparsity)	.13±.03	.86±.04	1.0±.00	.03±.01	.98±.02	1.0±.00
Random	–	.37±.02	.63±.02	–	.37±.02	.60±.02
Supervised	–	1.0±.00	1.0±.00	–	1.0±.00	1.0±.00

Datasets	TimeNonDiag _{NoSuffInf}			TimeNonDiag _{k=2}		
Metrics	SHD	MCC	R	SHD	MCC	R
PCL	–	.66±.04	.96±.00	–	.58±.04	.83±.01
SlowVAE (UDR)	–	.59±.02	.98±.00	–	.57±.01	.93±.00
SlowVAE (MCC)	–	.71±.02	.98±.00	–	.58±.02	.95±.00
TCVAE (UDR)	–	.58±.03	.98±.00	–	.57±.01	.96±.00
TCVAE (MCC)	–	.58±.03	.98±.00	–	.57±.01	.96±.00
Ours (no sparsity)	.45±.00	.62±.04	.98±.00	.45±.00	.62±.01	.98±.00
Ours (sparsity)	.32±.05	.63±.03	.99±.00	.20±.07	.74±.04	.98±.00
Random	–	.40±.04	.67±.02	–	.36±.01	.59±.02
Supervised	–	1.0±.00	1.0±.00	–	1.0±.00	1.0±.00

Table 3: Datasets ActionNonDiag_{NoSuffInf} and TimeNonDiag_{NoSuffInf} do not satisfy their respective sufficient influence assumptions (Assumptions 6 & 11). Datasets ActionNonDiag_{k=2} and TimeNonDiag_{k=2} are such that $\text{var}(\mathbf{z}^t \mid \mathbf{z}^{t-1}, \mathbf{a}^{t-1})$ depends on $\mathbf{z}^{t-1}$ or $\mathbf{a}^{t-1}$ (which means the sufficient statistic has dimension $k = 2$, contrarily to all other datasets). For our method, we show performance both with and without the sparsity constraint. In the former case, the constraint is set to the number of edges in the ground-truth graph. For baselines that have hyperparameters, we report their performance with the hyperparameter configurations that maximize UDR and MCC.

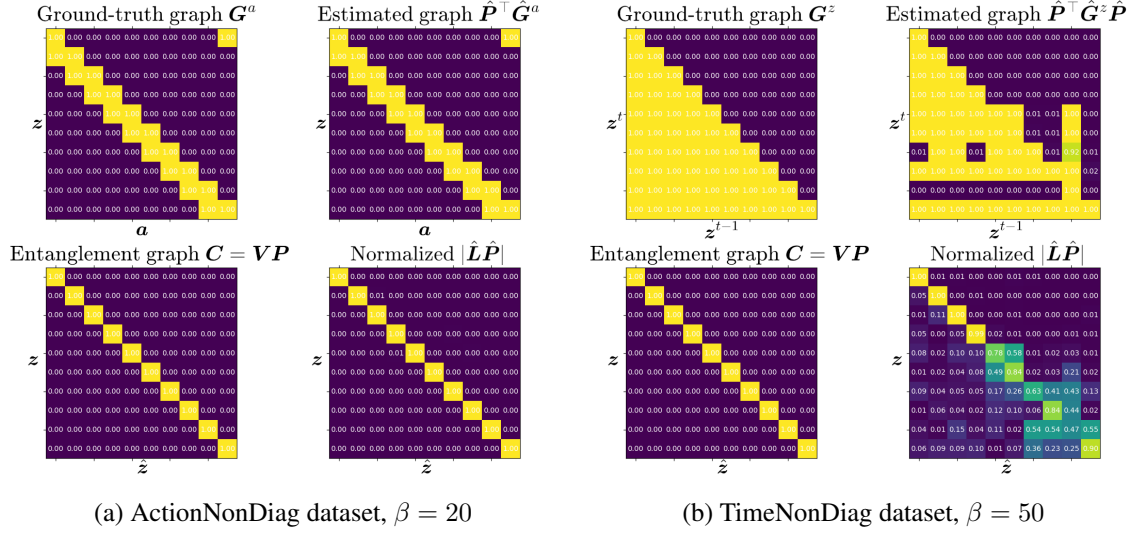


Figure 6: For each dataset, we visualize the median SHD run among the five randomly initialized runs of Figure 5 with the sparsity level β that is the closest to the ground-truth sparsity level $\|G\|_0$. For each dataset, we visualize (i) the ground-truth graph, (ii) the permuted estimated graph, (iii) the entanglement graph predicted by our theory, and (iv) the permuted matrix of regression coefficients in absolute value normalized by the maximum coefficient i.e. $|\hat{L}\hat{P}|/\max_{i,j} |\hat{L}_{i,j}|$. In Figure 6a, the estimated graph is exactly the ground-truth and $|\hat{L}\hat{P}|$ is perfectly diagonal, indicating complete disentanglement. In Figure 6b, the learned graph is close but not equal to the ground-truth. We can see that the off-diagonal nonzero values in $|\hat{L}\hat{P}|$ align with the poorly estimated parts of the graph.

8.2 Graphs Allowing only Partial Disentanglement (Violating Assumption 5)

In this section, we explore datasets with graphs that do not satisfy the criterion of Assumption 5. This means that our theory can only guarantee a form of partial disentanglement. For this reason, we will report the R_{con} metric introduced in Section 6 which measures whether two representations are a -consistent (Definition 13) or z -consistent (Definition 14).

Satisfying sufficient influence assumptions. Figure 7 reports the MCC, R_{con} and R scores of all methods on four datasets that satisfy the sufficient influence assumption but not the graphical criterion: the datasets ActionBlockDiag and TimeBlockDiag have “block diagonal” graphs, while ActionBlockNonDiag and TimeBlockNonDiag have more intricate graphs (depicted in Figure 8). See Appendix C.1 for details on the datasets. **Observations:** In all four datasets, some sparsity level yields near perfect R_{con} , indicating the learned models are approximately a -consistent or z -consistent to the ground-truth. Moreover, SHD is correlated with R_{con} . Without surprise, MCC never comes close to one, since complete disentanglement is not guaranteed by our theory. Similarly to Figure 5, the baselines have decent R values but very low MCC and R_{con} , indicating that they cannot achieve partial disentanglement. Figure 8 shows examples of estimated graphs. When it comes to hyperparameter selection, UDR selects the hyperparameter with the lowest SHD on three out of four datasets, which indicates that UDR does reasonably well.

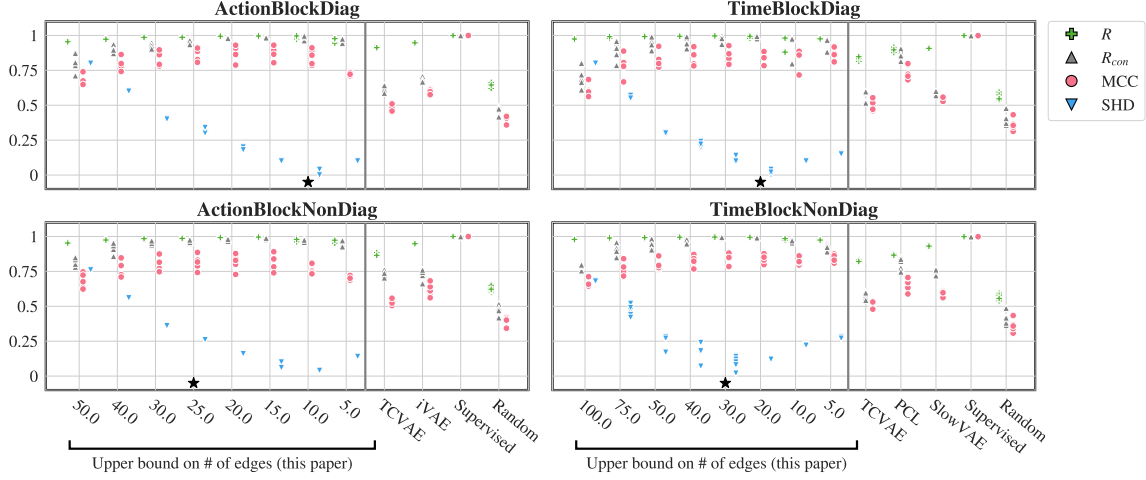


Figure 7: **Graphical criterion does not hold:** Datasets ActionBlockDiag and TimeBlockDiag have block-diagonal graphs while ActionBlockNonDiag and TimeBlockNonDiag have non-diagonal graphs. Sufficient influence is always satisfied. In the left column, only $\hat{G}^a$ is learned and we vary β_a , and in the right column, only $\hat{G}^z$ is learned and we vary β_z . For more details on the synthetic datasets, see Appendix C.1. The black star indicates which regularization parameter is selected by the filtered UDR procedure (see Appendix C.4). For R and MCC, higher is better. For SHD, lower is better. Performance is reported on 5 random seeds.

ActionRandomGraphs								
$p(\text{edge})$	Without sparsity			With sparsity				$\mathbb{E} V _0$
	MCC	R_{con}	R	MCC	R_{con}	R	SHD	
10%	.58±.13	.61±.11	.68±.13	.69±.14	.70±.12	.70±.12	.00±.00	45.0
20%	.67±.06	.69±.05	.83±.08	.85±.08	.86±.08	.86±.09	.01±.01	25.8
40%	.67±.03	.70±.03	.93±.04	.94±.05	.95±.05	.98±.04	.06±.05	15.8
60%	.69±.06	.73±.05	.96±.00	.88±.07	.91±.05	.99±.01	.14±.08	15.8
90%	.63±.04	.81±.08	.97±.00	.60±.01	.78±.07	.97±.00	.21±.07	45.0

TimeRandomGraphs								
$p(\text{edge})$	Without sparsity			With sparsity				$\mathbb{E} V _0$
	MCC	R_{con}	R	MCC	R_{con}	R	SHD	
10%	.66±.03	.66±.03	.98±.00	1.0±.00	1.0±.00	1.0±.00	.01±.02	10.2
20%	.63±.05	.63±.05	.98±.00	.99±.01	.99±.01	.99±.00	.08±.09	10.2
40%	.61±.02	.61±.02	.98±.00	.82±.16	.82±.16	.98±.01	.27±.13	10.2
60%	.58±.02	.58±.02	.98±.00	.71±.12	.71±.12	.98±.00	.33±.06	10.4
90%	.58±.03	.63±.08	.98±.00	.58±.02	.63±.08	.98±.00	.20±.07	26.1

Table 4: Experiments with randomly generated graphs. The probability of sampling an edge is $p(\text{edge})$. We report an estimation of $\mathbb{E}||V||_0$ which is the average number of edges in the entanglement graph V entailed by the random ground-truth graph G .

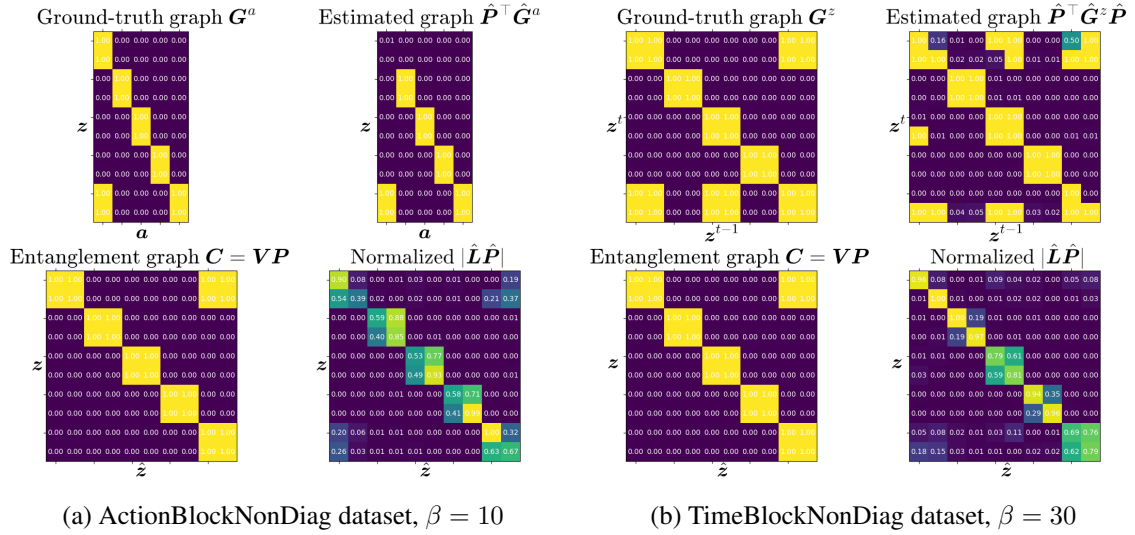


Figure 8: For each dataset, we visualize the median SHD run among the five randomly initialized runs of Figure 7 with the sparsity level β that is the closest to the ground-truth sparsity level $\|G\|_0$. For each dataset, we visualize (i) the ground-truth graph, (ii) the permuted estimated graph, (iii) the entanglement graph predicted by our theory, and (iv) the permuted matrix of regression coefficients in absolute value normalized by the maximum coefficient i.e. $|\hat{L}\hat{P}|/\max_{i,j} |\hat{L}_{i,j}|$. For both datasets, the learned graph is very close to the ground-truth. Furthermore, the match between the zero entries of $\hat{L}\hat{P}$ and those of the theoretical entanglement graph C is very good, although not perfect. Notice how certain blocks of latent factors remain entangled, as predicted by the theory.

Random graphs of varying sparsity levels. In Table 4, we consider the same μ functions as in datasets ActionNonDiag and TimeNonDiag, but explore more diverse randomly generated ground-truth graphs with various degrees of sparsity. Edges are sampled i.i.d. with some probability $p(\text{edge})$. However, note that for the TimeRandomGraphs dataset, the self-loops are present with probability one. We report the performance of our approach both with and without sparsity regularization. When using sparsity, we set β equal to the ground-truth number of edges, $\|\mathbf{G}\|_0$. **Observations:** First, all datasets obtain an improvement in MCC and R_{con} from sparsity regularization, except for very dense graphs with $p(\text{edge}) = 90\%$, in which case regularization does nothing or slightly degrades performance. Secondly, we can see that the SHD tends to be higher for larger graphs, suggesting that these are harder to learn. Third, in the ActionRandomGraphs datasets, we can see a negative correlation between MCC and $\mathbb{E}\|\mathbf{V}\|_0$, which is expected since $\|\mathbf{V}\|_0$ close to 10 means complete disentanglement is possible (assuming the graph is learned properly). This pattern also appears to some extent in the TimeRandomGraphs datasets. Notice how, among the TimeRandomGraphs datasets, all datasets sparser than $p(\text{edge}) = 90\%$ always have $\mathbb{E}\|\mathbf{V}\|_0 \approx 10$, indicating complete disentanglement should be possible.⁹ This is confirmed by very high MCC, at least for sparser graphs which are learned properly. Finally, we note that the R score is low for the very sparse action datasets. We suspect this is because very sparse graphs are less likely to satisfy the assumption of sufficient variability (Theorem 4), which guarantees quasi-linear equivalence (here it is actually *linear* equivalence, because of Gaussianity). Indeed, for very sparse graphs, some latent factors might end up without parents. This is not the case in the time datasets because of the self-loops which are always present.

9 Conclusion

This work proposed a novel principle for disentanglement based on *mechanism sparsity regularization*. The idea is based on the assumption that the mechanisms that govern the dynamics of high-level concepts are often sparse: actions usually affect only a few entities, and objects usually interact sparsely with each other. We provided novel nonparametric identifiability guarantees for this setting which give sufficient conditions for disentanglement, whether complete or partial. Given the dependency structure between latent factors and auxiliary variables, our theory predicts the entanglement graph that describes which of the estimated latent factors are expected to remain entangled. This constitutes a significant extension of the shorter version of this work (Lachapelle et al., 2022). We further provide various examples to illustrate the consequences of our guarantees as well as the assumptions they rely on. For instance, we show that multi-node interventions with unknown targets fall under the umbrella of our framework. Finally, we demonstrate the theory experimentally by training a sparsity-constrained variational autoencoder on synthetic data, which allows us to explore various settings. Our work establishes a solid theoretical foundation for further empirical investigations in more realistic scenarios, such as single-cell data with gene perturbations (Lopez et al., 2023) and video (Lei et al., 2023). Future works include relaxing assumptions such as conditional independence or considering more permissive settings such as “contextual sparsity”, i.e., the assumption that objects only interact with each other in particular situations. We believe that the latter could be formalized and leveraged for disentanglement using the tools developed in this work.

9. We suspect this occurs because the self-loops which are present with probability one, unlike the Action dataset.

Acknowledgments and Disclosure of Funding

The authors would like to thank Aristide Baratin for important feedback on the manuscript. This research was partially supported by the Canada CIFAR AI Chair Program, by an IVADO excellence PhD scholarship, by a Google Focused Research award, the German Federal Ministry of Education and Research (BMBF): Tübingen AI Center, FKZ: 01IS18039A, and by Mitacs through the Mitacs Accelerate program. The experiments were in part enabled by computational resources provided by Calcul Quebec and the Digital Research Alliance of Canada. The authors would like to thank Yoshua Bengio for inspiring mechanism sparsity regularization through various talks and discussions. The authors would also like to thank the International Max Planck Research School for Intelligent Systems (IMPRS-IS) for supporting Yash Sharma. Simon Lacoste-Julien is a CIFAR Associate Fellow in the Learning in Machines & Brains program.

<u>Calligraphic & indexing conventions</u>	
$[n]$	$\{1, 2, \dots, n\}$
x	Scalar (random or not, depending on context)
$\mathbf{x}$	Vector (random or not, depending on context)
$\mathbf{X}$	Matrix
$\mathcal{X}$	Set/Support
f	Scalar-valued function
$\mathbf{f}$	Vector-valued function
$Df, D\mathbf{f}$	Jacobian of f and $\mathbf{f}$
D^2f	Hessian of f
$B \subseteq [n]$	Subset of indices
$\mathbf{x}_B$	Vector formed with the i th coordinates of $\mathbf{x}$, for all $i \in B$
$\mathbf{X}_{B,B'}$	Matrix formed with the entries $(i, j) \in B \times B'$ of $\mathbf{X}$.
 <u>Recurrent notation</u>	
$\mathbf{x}^t \in \mathbb{R}^{d_x}$	Observation at time t
$\mathbf{x}^{\leq t} \in \mathbb{R}^{d_x \times t}$	Matrix of observations at times $1, \dots, t$
$\mathbf{z}^t \in \mathbb{R}^{d_z}$	Vector of latent factors of variations at time t
$\mathbf{z}^{\leq t} \in \mathbb{R}^{d_z \times t}$	Matrix of latent vectors at times $1, \dots, t$
$\mathbf{a}^t \in \mathbb{R}^{d_a}$	Vector of auxiliary variables at time t
$\mathbf{a}^{< t} \in \mathbb{R}^{d_a \times t}$	Matrix of auxiliary vectors at times $0, 1, \dots, t-1$
$\mathcal{A} \subseteq \mathbb{R}^{d_a}$	Support of $\mathbf{a}^t$
$\mathbf{f} : \mathbb{R}^{d_z} \rightarrow \mathbb{R}^{d_x}$	Ground-truth decoder function
$\hat{\mathbf{f}} : \mathbb{R}^{d_z} \rightarrow \mathbb{R}^{d_x}$	Learned decoder function
$p(\mathbf{z}^t \mid \mathbf{z}^{< t}, \mathbf{a}^{< t})$	Ground-truth latent transition model
$\hat{p}(\mathbf{z}^t \mid \mathbf{z}^{< t}, \mathbf{a}^{< t})$	Learned latent transition model
$\mathbf{G}^a \in \{0, 1\}^{d_z \times d_a}$	Ground-truth adjacency matrix of graph connecting $\mathbf{a}^{< t}$ to $\mathbf{z}^t$
$\mathbf{G}^z \in \{0, 1\}^{d_z \times d_z}$	Ground-truth adjacency matrix of graph connecting $\mathbf{z}^{< t}$ to $\mathbf{z}^t$
$\hat{\mathbf{G}}^a, \hat{\mathbf{G}}^z$	Learned adjacency matrices
$\mathbf{Pa}_i^a \subseteq [d_a]$	Parents of $\mathbf{z}_i^t$ in $\mathbf{G}^a$
$\mathbf{Ch}_\ell^a \subseteq [d_z]$	Children of $\mathbf{a}_\ell^t$ in $\mathbf{G}^z$
$\mathbf{Pa}_i^z \subseteq [d_z]$	Parents of $\mathbf{z}_i^t$ in $\mathbf{G}^z$
$\mathbf{Ch}_i^z \subseteq [d_z]$	Children of $\mathbf{z}_i^{t-1}$ in $\mathbf{G}^z$
$D_z^t \log p \in \mathbb{R}^{1 \times d_z}$	Jacobian vector of $\log p(\mathbf{z}^t \mid \mathbf{z}^{< t}, \mathbf{a}^{< t})$ w.r.t. $\mathbf{z}^t$
$H_{z,a}^{t,\tau} \log p \in \mathbb{R}^{d_z \times d_a}$	Hessian matrix of $\log p(\mathbf{z}^t \mid \mathbf{z}^{< t}, \mathbf{a}^{< t})$ w.r.t. $\mathbf{z}^t$ and $\mathbf{a}^\tau$
$H_{z,z}^{t,\tau} \log p \in \mathbb{R}^{d_z \times d_z}$	Hessian matrix of $\log p(\mathbf{z}^t \mid \mathbf{z}^{< t}, \mathbf{a}^{< t})$ w.r.t. $\mathbf{z}^t$ and $\mathbf{z}^\tau$
$\sigma : [d_z] \rightarrow [d_z]$	A permutation
 <u>Topology</u>	
$\overline{\mathcal{X}}$	Closure of the set $\mathcal{X} \subseteq \mathbb{R}^n$
$\mathcal{X}^\circ$	Interior of the set $\mathcal{X} \subseteq \mathbb{R}^n$

Table 5: Table of Notation.

Appendix A. Identifiability Theory - Nonparametric Case

A.1 Useful Lemmas

Definition 18 (Regular closed set) A set $A \subseteq \mathbb{R}^n$ is regular closed when it is equal to the closure of its interior, i.e. $\overline{A^\circ} = A$.

Lemma 4 Let $A \subseteq \mathbb{R}^n$ and $\mathbf{f} : A \rightarrow \mathbb{R}^m$ be a C^k function. Then, its k first derivatives is uniquely defined on $\overline{A^\circ}$ in the sense that they do not depend on the specific choice of C^k extension.

Proof Let $\mathbf{g} : U \rightarrow \mathbb{R}^m$ and $\mathbf{h} : V \rightarrow \mathbb{R}^m$ be two C^k extensions of $\mathbf{f}$ to $U \subseteq \mathbb{R}^n$ and $V \subseteq \mathbb{R}^n$ both open in $\mathbb{R}^n$. By definition,

$$\mathbf{g}(\mathbf{x}) = \mathbf{f}(\mathbf{x}) = \mathbf{h}(\mathbf{x}), \forall \mathbf{x} \in A.$$

The usual derivative is uniquely defined on the interior of the domain, so that

$$D\mathbf{g}(\mathbf{x}) = D\mathbf{f}(\mathbf{x}) = D\mathbf{h}(\mathbf{x}), \forall \mathbf{x} \in A^\circ.$$

Consider a point $\mathbf{x}_0 \in \overline{A^\circ}$. By definition of closure, there exists a sequence $\{\mathbf{x}_k\}_{k=1}^\infty \subseteq A^\circ$ s.t. $\lim_{k \rightarrow \infty} \mathbf{x}_k = \mathbf{x}_0$. We thus have that

$$\begin{aligned} \lim_{k \rightarrow \infty} D\mathbf{g}(\mathbf{x}_k) &= \lim_{k \rightarrow \infty} D\mathbf{h}(\mathbf{x}_k) \\ D\mathbf{g}(\mathbf{x}_0) &= D\mathbf{h}(\mathbf{x}_0), \end{aligned}$$

where we used the fact that the derivatives of $\mathbf{g}$ and $\mathbf{h}$ are continuous to go to the second line. Thus, all the C^k extensions of $\mathbf{f}$ must have equal derivatives on $\overline{A^\circ}$. This means we can unambiguously define the derivative of $\mathbf{f}$ everywhere on $\overline{A^\circ}$ to be equal to the derivative of one of its C^k extensions.

Since $\mathbf{f}$ is C^k , its derivative $D\mathbf{f}$ is C^{k-1} , we can thus apply the same argument to get that the second derivative of $\mathbf{f}$ is uniquely defined on $\overline{A^\circ}$. It can be shown that $\overline{A^{\circ\circ}} = \overline{A^\circ}$. One can thus apply the same argument recursively to show that the first k derivatives of $\mathbf{f}$ are uniquely defined on $\overline{A^\circ}$. ■

Lemma 5 Let X be some set. A family of functions $(f_i : X \rightarrow \mathbb{R})_{i=1}^n$ is linearly independent if and only if there exists $x_1, \dots, x_n \in X$ such that the family of vectors $((f_1(x_i), \dots, f_n(x_i)))_{i=1}^n$ is linearly independent.

Proof We start by proving the “if” part. Assume the functions are linearly dependent. Then there exists $\alpha \neq 0$ such that, for all $x \in X$, $\sum_{i=1}^n \alpha_i f_i(x) = 0$. Choose distinct $x_1, \dots, x_n \in X$. We thus have that for all $j \in [n]$, $\sum_{i=1}^n \alpha_i f_i(x_j) = 0$. This can be written in matrix form:

$$\begin{bmatrix} f_1(x_1) & \cdots & f_1(x_n) \\ \vdots & \ddots & \vdots \\ f_n(x_1) & \cdots & f_n(x_n) \end{bmatrix} \alpha = \mathbf{0},$$

which implies that the columns are linearly dependent.

We now show the “only if” part. Suppose that for all $\{x_1, \dots, x_n\} \subseteq X$, the family of vectors $((f_1(x_i), \dots, f_n(x_i)))_{i=1}^n$ is linearly dependent. This means that the set $U = \text{span}\{(f_1(x), \dots, f_n(x)) \mid x \in X\}$ is a proper linear subspace of $\mathbb{R}^n$. This means that there is a nonzero $u \in U^\perp$, the orthogonal complement of U . By definition, u is orthogonal to all elements in $\{(f_1(x), \dots, f_n(x)) \mid x \in X\}$. In other words, for all $x \in X$, $\sum_{i=1}^n u_i f_i(x) = 0$. Hence the f_i are linearly dependent. $\blacksquare$

A.2 Proof of Proposition 1

Proposition 1 (Linking dependency graph and Jacobian) *Let $\mathbf{h}$ be a C^1 function, i.e. continuously differentiable, from $\mathbb{R}^n$ to $\mathbb{R}^m$ and let $\mathbf{H}$ be its dependency graph (Definition 2). Then,*

$$\mathbf{H}_{i,j} = 0 \iff \text{For all } \mathbf{a} \in \mathbb{R}^n, D\mathbf{h}(\mathbf{a})_{i,j} = 0. \quad (4)$$

Proof The “ $\implies$ ” direction holds since we can simply differentiate $\mathbf{h}_i(\mathbf{a}) = \bar{\mathbf{h}}_i(\mathbf{a}_{-j})$ w.r.t. $\mathbf{a}_j$ to get zero.

We now show the “ $\impliedby$ ” direction. Suppose that for all $\mathbf{a} \in \mathbb{R}^n$, $D\mathbf{h}(\mathbf{a})_{i,j} = 0$. We must now show that $\mathbf{h}_i(\mathbf{a})$ is constant in $\mathbf{a}_j$ for all $\mathbf{a}_{-j}$. Choose any $\mathbf{a}^0, \mathbf{a}^1 \in \mathbb{R}^n$ such that $\mathbf{a}_{-j}^0 = \mathbf{a}_{-j}^1$. Thanks to the fundamental theorem of calculus, we can write

$$\begin{aligned} \mathbf{h}_i(\mathbf{a}^1) - \mathbf{h}_i(\mathbf{a}^0) &= \int_{[0,1]} \frac{d}{d\alpha} \mathbf{h}_i((1-\alpha)\mathbf{a}^0 + \alpha\mathbf{a}^1) d\alpha \\ &= \int_{[0,1]} \underbrace{D\mathbf{h}((1-\alpha)\mathbf{a}^0 + \alpha\mathbf{a}^1)_{i,\cdot}}_{\text{zero at } j} \cdot \underbrace{(\mathbf{a}^1 - \mathbf{a}^0)}_{\text{zero except at } j} d\alpha \\ &= 0. \end{aligned}$$

Since $\mathbf{a}^0$ and $\mathbf{a}^1$ were arbitrary points such that $\mathbf{a}_{-j}^0 = \mathbf{a}_{-j}^1$, this means the function $\mathbf{h}_i(\mathbf{a})$ is constant in $\mathbf{a}_j$ for all values of $\mathbf{a}_{-j}$. $\blacksquare$

A.3 Proof of Proposition 2

In this section, we prove Proposition 2. Before doing so, we first recall the definition of the support of a random variable (Definition 19) and prove a useful lemma (Lemma 6).

Definition 19 (Support of a random variable) *Let $\mathbf{x}$ be a random variable with values in $\mathbb{R}^n$ with distribution $\mathbb{P}_{\mathbf{x}}$. Let $\mathcal{O}_n$ be the standard topology of $\mathbb{R}^n$ (i.e. the set of open sets of $\mathbb{R}^n$). The support of $\mathbf{x}$ is defined as*

$$\text{supp}(\mathbf{x}) := \{\mathbf{x} \in \mathbb{R}^n \mid \mathbf{x} \in O \in \mathcal{O}_n \implies \mathbb{P}_{\mathbf{x}}(O) > 0\}.$$

Lemma 6 *Let $\mathbf{z}$ be a random variable with values in $\mathbb{R}^m$ with distribution $\mathbb{P}_{\mathbf{z}}$ and $\mathbf{y} := \mathbf{f}(\mathbf{z})$ where $\mathbf{f} : \text{supp}(\mathbf{z}) \rightarrow \mathbb{R}^n$ is a homeomorphism onto its image. Then*

$$\mathbf{f}(\text{supp}(\mathbf{z})) \subseteq \text{supp}(\mathbf{y}) \subseteq \overline{\mathbf{f}(\text{supp}(\mathbf{z}))}.$$

where the closure is taken w.r.t. the topology of $\mathbb{R}^n$.

Proof We first prove that $f(\text{supp}(z)) \subseteq \text{supp}(y)$. Let $y^0 \in f(\text{supp}(z))$ and N be an open neighborhood of y^0 , i.e. $y^0 \in N \in \mathcal{O}_n$. Note that there exists $z^0 \in \text{supp}(z)$ such that $f(z^0) = y^0$. Note that $z^0 \in f^{-1}(\{y^0\}) \subseteq f^{-1}(N)$ and that, by continuity of f , $f^{-1}(N)$ is an open neighborhood of z^0 . Since $z^0 \in \text{supp}(z)$, we have

$$\begin{aligned} 0 &< \mathbb{P}_z(f^{-1}(N)) \\ &= \mathbb{P}_z \circ f^{-1}(N) \\ &= \mathbb{P}_y(N). \end{aligned}$$

Hence $y^0 \in \text{supp}(y)$, which concludes the “ $\subseteq$ ” part.

We now prove the other inclusion. Let $y^0 \in \text{supp}(y)$ and suppose, by contradiction, that $y^0 \notin \overline{f(\text{supp}(z))}$. Since $\overline{f(\text{supp}(z))}$ is closed in $\mathbb{R}^n$, there exists N s.t. $y^0 \in N \in \mathcal{O}_n$ with $N \cap \overline{f(\text{supp}(z))} = \emptyset$. Since $y^0 \in \text{supp}(y)$,

$$\begin{aligned} 0 &< \mathbb{P}_y(N) \\ &= \mathbb{P}_z(f^{-1}(N)) \\ &= \mathbb{P}_z(\emptyset) = 0. \end{aligned}$$

The above contradiction implies that $y^0 \in \overline{f(\text{supp}(z))}$. ■

Proposition 2 (Identifiability up to diffeomorphism) *Let $\theta := (f, p, G)$ and $\hat{\theta} := (\hat{f}, \hat{p}, \hat{G})$ be two models satisfying Assumption 1. If $\theta \sim_{\text{obs}} \hat{\theta}$ (Def. 4), then $\theta \sim_{\text{diff}} \hat{\theta}$ (Def. 5).*

Proof

Equality of Denoised Distributions. Given an arbitrary $a^{<T} \in \mathcal{A}^T$ and a parameter $\theta = (f, p, G)$, let $\mathbb{P}_{x \leq T | a^{<T}; \theta}$ be the conditional probability distribution of $x \leq T$, let $\mathbb{P}_{z \leq T | a^{<T}; \theta}$ be the conditional probability distribution of $z \leq T$ and let $\mathbb{P}_{n \leq T}$ be the probability distribution of $n \leq T$ (the Gaussian noises added on $f(z \leq T)$, defined in Sec. 2.1). Let $y^t := f(z^t)$ and $\mathbb{P}_{y \leq T | a^{<T}; \theta}$ be its conditional probability distribution. First, notice that

$$\mathbb{P}_{x \leq T | a^{<T}; \theta} = \mathbb{P}_{y \leq T | a^{<T}; \theta} * \mathbb{P}_{n \leq T},$$

where $*$ is the convolution operator between two measures. We now show that if two models agree on the observations, i.e. $\mathbb{P}_{x \leq T | a^{<T}; \theta} = \mathbb{P}_{x \leq T | a^{<T}; \hat{\theta}}$, then $\mathbb{P}_{y \leq T | a^{<T}; \theta} = \mathbb{P}_{y \leq T | a^{<T}; \hat{\theta}}$. The following argument makes use of the Fourier transform $\mathcal{F}$ generalized to arbitrary probability measures. This tool is necessary to deal with measures which do not have a density w.r.t either the Lebesgue or the counting measure, as is the case of $\mathbb{P}_{y \leq T | a^{<T}; \theta}$ (all its mass is concentrated on the set $f(\mathbb{R}^{d_z})$). See Pollard (2001, Chapter 8) for an introduction and useful properties.

$$\mathbb{P}_{x \leq T | a^{<T}; \theta} = \mathbb{P}_{x \leq T | a^{<T}; \hat{\theta}} \tag{34}$$

$$\mathbb{P}_{y \leq T | a^{<T}; \theta} * \mathbb{P}_{n \leq T} = \mathbb{P}_{y \leq T | a^{<T}; \hat{\theta}} * \mathbb{P}_{n \leq T} \tag{35}$$

$$\mathcal{F}(\mathbb{P}_{y \leq T | a^{<T}; \theta} * \mathbb{P}_{n \leq T}) = \mathcal{F}(\mathbb{P}_{y \leq T | a^{<T}; \hat{\theta}} * \mathbb{P}_{n \leq T}) \tag{36}$$

$$\mathcal{F}(\mathbb{P}_{y \leq T | a^{<T}; \theta}) \mathcal{F}(\mathbb{P}_{n \leq T}) = \mathcal{F}(\mathbb{P}_{y \leq T | a^{<T}; \hat{\theta}}) \mathcal{F}(\mathbb{P}_{n \leq T}) \tag{37}$$

$$\mathcal{F}(\mathbb{P}_{y \leq T | a^{<T}; \theta}) = \mathcal{F}(\mathbb{P}_{y \leq T | a^{<T}; \hat{\theta}}) \tag{38}$$

$$\mathbb{P}_{y \leq T | a^{<T}; \theta} = \mathbb{P}_{y \leq T | a^{<T}; \hat{\theta}}, \tag{39}$$

where (36) & (39) use the fact that the Fourier transform is invertible, (37) is an application of the fact that the Fourier transform of a convolution is the product of their Fourier transforms and (38) holds because the Fourier transform of a Normal distribution is nonzero everywhere. Note that the latter argument holds because we assume σ^2 , the variance of the Gaussian noise added to $\mathbf{y}^t$, is the same for both models. Notice that, since $\mathbf{f}(\mathbb{R}^{d_z})$ and $\hat{\mathbf{f}}(\mathbb{R}^{d_z})$ are closed in $\mathbb{R}^{d_x}$ and the support of $\mathbf{z}^{\leq T}$ is $\mathbb{R}^{d_z \times T}$, Lemma 6 implies that

$$\mathbf{f}(\mathbb{R}^{d_z \times T}) = \text{supp}(\mathbb{P}_{\mathbf{y}^{\leq T} | \mathbf{a}^{< T}; \boldsymbol{\theta}}) \quad \& \quad \text{supp}(\mathbb{P}_{\mathbf{y}^{\leq T} | \mathbf{a}^{< T}; \hat{\boldsymbol{\theta}}}) = \hat{\mathbf{f}}(\mathbb{R}^{d_z \times T})$$

where we overloaded the notation by defining $\mathbf{f}(\mathbf{z}^{\leq T}) := (\mathbf{f}(\mathbf{z}^1), \dots, \mathbf{f}(\mathbf{z}^T))$ and analogously for $\hat{\mathbf{f}}(\mathbf{z}^{\leq T})$. Since both measure in (39) are equal, their supports must also be. This implies that $\mathbf{f}(\mathbb{R}^{d_z}) = \hat{\mathbf{f}}(\mathbb{R}^{d_z})$, which is part of the definition of equivalence up to diffeomorphism (Definition 5).

Equality of densities. Continuing with (39),

$$\begin{aligned} \mathbb{P}_{\mathbf{y}^{\leq T} | \mathbf{a}^{< T}; \boldsymbol{\theta}} &= \mathbb{P}_{\mathbf{y}^{\leq T} | \mathbf{a}^{< T}; \hat{\boldsymbol{\theta}}} \\ \mathbb{P}_{\mathbf{z}^{\leq T} | \mathbf{a}^{< T}; \boldsymbol{\theta}} \circ \mathbf{f}^{-1} &= \mathbb{P}_{\mathbf{z}^{\leq T} | \mathbf{a}^{< T}; \hat{\boldsymbol{\theta}}} \circ \hat{\mathbf{f}}^{-1} \\ \mathbb{P}_{\mathbf{z}^{\leq T} | \mathbf{a}^{< T}; \boldsymbol{\theta}} \circ \mathbf{f}^{-1} \circ \hat{\mathbf{f}} &= \mathbb{P}_{\mathbf{z}^{\leq T} | \mathbf{a}^{< T}; \hat{\boldsymbol{\theta}}} \\ \mathbb{P}_{\mathbf{z}^{\leq T} | \mathbf{a}^{< T}; \boldsymbol{\theta}} \circ \mathbf{v} &= \mathbb{P}_{\mathbf{z}^{\leq T} | \mathbf{a}^{< T}; \hat{\boldsymbol{\theta}}}, \end{aligned} \tag{40}$$

where $\mathbf{v} := \mathbf{f}^{-1} \circ \hat{\mathbf{f}}$ is a composition of diffeomorphisms and thus a diffeomorphism from $\mathbb{R}^{d_z}$ to itself. Note that this composition is well defined because $\mathbf{f}(\mathbb{R}^{d_z}) = \hat{\mathbf{f}}(\mathbb{R}^{d_z})$. We chose to work directly with measures (functions on sets), as opposed to manifold integrals in Khemakhem et al. (2020a), because it simplifies the derivation of (40) and avoids having to define densities w.r.t. measures concentrated on a manifold.

The density of $\mathbb{P}_{\mathbf{z}^{\leq T} | \mathbf{a}^{< T}; \boldsymbol{\theta}} \circ \mathbf{v}$ w.r.t. to the Lebesgue measure is given by the change-of-variable rule for random vectors (which can be applied because $\mathbf{v}$ is a diffeomorphism) and is given by $\prod_{t=1}^T p(\mathbf{v}(\mathbf{z}^t) | \mathbf{v}(\mathbf{z}^{< t}), \mathbf{a}^{< t}) |\det D\mathbf{v}(\mathbf{z}^t)|$, where p refers to the density model with parameter $\boldsymbol{\theta}$ and $D\mathbf{v}(\mathbf{z}^t)$ is the Jacobian matrix of $\mathbf{v}$. Since $\mathbb{P}_{\mathbf{z}^{\leq T} | \mathbf{a}^{< T}; \boldsymbol{\theta}} \circ \mathbf{v} = \mathbb{P}_{\mathbf{z}^{\leq T} | \mathbf{a}^{< T}; \hat{\boldsymbol{\theta}}}$, their respective densities w.r.t. Lebesgue must also agree:

$$\prod_{t=1}^T \hat{p}(\mathbf{z}^t | \mathbf{z}^{< t}, \mathbf{a}^{< t}) = \prod_{t=1}^T p(\mathbf{v}(\mathbf{z}^t) | \mathbf{v}(\mathbf{z}^{< t}), \mathbf{a}^{< t}) |\det D\mathbf{v}(\mathbf{z}^t)|,$$

where $\hat{p}$ refers to the conditional density of the model with parameter $\hat{\boldsymbol{\theta}}$.

For a given t_0 , we have

$$\prod_{t=1}^{t_0} \hat{p}(\mathbf{z}^t | \mathbf{z}^{< t}, \mathbf{a}^{< t}) = \prod_{t=1}^{t_0} p(\mathbf{v}(\mathbf{z}^t) | \mathbf{v}(\mathbf{z}^{< t}), \mathbf{a}^{< t}) |\det D\mathbf{v}(\mathbf{z}^t)|, \tag{41}$$

by integrating first $\mathbf{z}^T$, then $\mathbf{z}^{t-1}$, then ..., up to $\mathbf{z}^{t_0+1}$. Note that we can integrate $\mathbf{z}^{t_0}$ and get

$$\prod_{t=1}^{t_0-1} \hat{p}(\mathbf{z}^t | \mathbf{z}^{< t}, \mathbf{a}^{< t}) = \prod_{t=1}^{t_0-1} p(\mathbf{v}(\mathbf{z}^t) | \mathbf{v}(\mathbf{z}^{< t}), \mathbf{a}^{< t}) |\det D\mathbf{v}(\mathbf{z}^t)|. \tag{42}$$

By dividing (41) by (42), we get

$$\hat{p}(\mathbf{z}^{t_0} | \mathbf{z}^{<t_0}, \mathbf{a}^{<t_0}) = p(\mathbf{v}(\mathbf{z}^{t_0}) | \mathbf{v}(\mathbf{z}^{<t_0}), \mathbf{a}^{<t_0}) |\det D\mathbf{v}(\mathbf{z}^{t_0})|, \quad (43)$$

which completes the proof. $\blacksquare$

A.4 The Consistency Relations (Definitions 13 & 14) are Equivalence Relations

In this section, we demonstrate that the relations $\sim_{\text{con}}^a$ and $\sim_{\text{con}}^z$ are equivalence relations by leveraging the fact that the set of G -preserving diffeomorphisms form a group under composition (Proposition 5). We start by showing a fact that will be useful in the following.

Lemma 7 *Let $G \in \{0, 1\}^{m \times n}$.*

1. *A map $c : \mathbb{R}^m \rightarrow \mathbb{R}^m$ is G -preserving if and only if c is GP -preserving, where P is an $n \times n$ permutation matrix.*
2. *A map $c : \mathbb{R}^m \rightarrow \mathbb{R}^m$ is G -preserving if and only if $P \circ c \circ P^\top$ is PG -preserving, where P is a $m \times m$ permutation matrix.*
3. *When $m = n$, a map $c : \mathbb{R}^m \rightarrow \mathbb{R}^m$ is G -preserving if and only if $P \circ c \circ P^\top$ is $PGP^\top$ -preserving, where P is a $m \times m$ permutation matrix.*

Proof Let C be the dependency graph of c .

1. $C^\top \mathbb{R}_G^{m \times n} \subseteq \mathbb{R}_G^{m \times n} \iff C^\top \mathbb{R}_G^{m \times n} P \subseteq \mathbb{R}_G^{m \times n} P \iff C^\top \mathbb{R}_{GP}^{m \times n} \subseteq \mathbb{R}_{GP}^{m \times n}$
2. First, notice that the dependency graph of $P \circ c \circ P^\top$ is $PCP^\top$.
 $(PCP^\top)^\top \mathbb{R}_{PG}^{m \times n} \subseteq \mathbb{R}_{PG}^{m \times n} \iff PC^\top P^\top P \mathbb{R}_G^{m \times n} \subseteq P \mathbb{R}_G^{m \times n} \iff C^\top \mathbb{R}_G^{m \times n} \subseteq \mathbb{R}_G^{m \times n}$
3. This is a consequence of the first two statements. $\blacksquare$

We are now ready to show that the relation $\sim_{\text{con}}^a$ (Definition 13) is an equivalence relation.

Proposition 9 *The consistency relation, $\sim_{\text{con}}^a$ (Def. 13), is an equivalence relation.*

Proof

Reflexivity. It is easy to see that $\theta \sim_{\text{con}}^a \theta$, simply by setting $v(z) := z$ with $P := I$.

Symmetry. Assume $\theta \sim_{\text{con}}^a \tilde{\theta}$. Hence, we have $PG^a = \tilde{G}^a$ as well as

$$f(\mathbb{R}^{d_z}) = \tilde{f}(\mathbb{R}^{d_z}), \text{ and} \quad (44)$$

$$\tilde{p}(\mathbf{z}^t | \mathbf{z}^{<t}, \mathbf{a}^{<t}) = p(\mathbf{v}(\mathbf{z}^t) | \mathbf{v}(\mathbf{z}^{<t}), \mathbf{a}^{<t}) |\det D\mathbf{v}(\mathbf{z}^t)|, \quad (45)$$

where $v := f^{-1} \circ \tilde{f}$ can be written as $v := c \circ P^\top$, where c is a G^a -preserving diffeomorphism and P is a permutation. We can massage (45) to get

$$p(\mathbf{z}^t | \mathbf{z}^{<t}, \mathbf{a}^{<t}) = \tilde{p}(v^{-1}(\mathbf{z}^t) | v^{-1}(\mathbf{z}^{<t}), \mathbf{a}^{<t}) |\det Dv^{-1}(\mathbf{z}^t)|.$$

Of course, we also have that $\tilde{P}\tilde{G}^a = G^a$, where $\tilde{P} := P^\top$. Now the only thing left to prove is that v^{-1} can be written as $\tilde{c} \circ \tilde{P}^\top$ where $\tilde{c}$ is $\tilde{G}^a$ -preserving. We know that

$$v^{-1} = P \circ c^{-1} = \underbrace{P \circ c^{-1} \circ P^\top}_{\tilde{c}:=} \circ P = \tilde{c} \circ \tilde{P}^\top.$$

Note that c^{-1} is G^a -preserving and thus, by Lemma 7, $\tilde{c}$ is PG^a -preserving, i.e. $\tilde{G}^a$ -preserving. Hence, $\sim_{\text{con}}^a$ is symmetric.

Transitivity. Suppose $\theta \sim_{\text{con}}^a \tilde{\theta}$ and $\tilde{\theta} \sim_{\text{con}}^a \hat{\theta}$. This means

$$P_1 G^a = \tilde{G}^a, \quad (46)$$

$$f(\mathbb{R}^{d_z}) = \tilde{f}(\mathbb{R}^{d_z}), \text{ and} \quad (47)$$

$$\tilde{p}(z^t \mid z^{<t}, a^{<t}) = p(v_1(z^t) \mid v_1(z^{<t}), a^{<t}) |\det Dv_1(z^t)|, \quad (48)$$

where $v_1 := c_1 \circ P_1^\top$ with c_1 being G^a -preserving; and

$$P_2 \tilde{G}^a = \hat{G}^a, \quad (49)$$

$$\tilde{f}(\mathbb{R}^{d_z}) = \hat{f}(\mathbb{R}^{d_z}), \text{ and} \quad (50)$$

$$\hat{p}(z^t \mid z^{<t}, a^{<t}) = \tilde{p}(v_2(z^t) \mid v_2(z^{<t}), a^{<t}) |\det Dv_2(z^t)|, \quad (51)$$

where $v_2 := c_2 \circ P_2^\top$ with c_2 being $\tilde{G}^a$ -preserving.

To show that $\theta \sim_{\text{con}}^a \hat{\theta}$, we first combine (46) with (49) to get

$$\underbrace{P_2 P_1}_{P:=} G^a = \hat{G}^a. \quad (52)$$

Of course we also have that $f(\mathbb{R}^{d_z}) = \tilde{f}(\mathbb{R}^{d_z}) = \hat{f}(\mathbb{R}^{d_z})$. By massaging both (48) and (51), we get:

$$\hat{p}(z^t \mid z^{<t}, a^{<t}) = p(v_1 \circ v_2(z^t) \mid v_1 \circ v_2(z^{<t}), a^{<t}) |\det D(v_1 \circ v_2)(z^t)|.$$

Define $v := v_1 \circ v_2$. We now want to show that v can be written as $v = c \circ P^\top$ where c is G^a -preserving. We have that

$$\begin{aligned} v_1 \circ v_2 &= c_1 \circ P_1^\top \circ c_2 \circ P_2^\top \\ &= c_1 \circ \underbrace{P_1^\top \circ P_2^\top}_{P^\top=} \circ \underbrace{P_2 \circ c_2 \circ P_2^\top}_{\hat{c}:=} \\ &= c_1 \circ P^\top \circ \hat{c} \end{aligned}$$

where, by Lemma 7, $\hat{c}$ is $P_2 \tilde{G}^a$ -preserving, i.e. $\hat{G}^a$ -preserving. We continue and get

$$\begin{aligned} v_1 \circ v_2 &= c_1 \circ P^\top \circ \hat{c} \\ &= c_1 \circ \underbrace{P^\top \circ \hat{c} \circ P}_{c':=} \circ P^\top \\ &= c_1 \circ c' \circ P^\top, \end{aligned}$$

where, by Lemma 7, c' is $P^\top \hat{G}^a$ -preserving, i.e. G^a -preserving (by (52)). Since both c_1 and c' are G^a -preserving, $c := c_1 \circ c'$ is G^a -preserving, which concludes the proof. $\blacksquare$

The same can be shown for $\sim_{\text{con}}^z$ (Definition 14).

Proposition 10 *The consistency relation, $\sim_{\text{con}}^z$ (Def. 14), is an equivalence relation.*

Proof The proof is exactly analogous to the proof that $\sim_{\text{con}}^a$ is an equivalence relation. Essentially, every statement of the form “ $PG^a = \tilde{G}^a$ ” becomes “ $PG^z P^\top = \tilde{G}^z$ ” and statements of the form “ c is G^a -preserving” becomes “ c is G^z -preserving and $(G^z)^\top$ -preserving”. The full proof is left as an exercise to the reader. $\blacksquare$

A.4.1 COMBINING EQUIVALENCE RELATIONS

Proposition 6 *Let $\theta := (f, p, G)$ and $\tilde{\theta} := (\tilde{f}, \tilde{p}, \tilde{G})$ be two models satisfying Assumptions 1, 2 & 3. We have $\theta \sim_{\text{con}}^z \tilde{\theta}$ if and only if $\theta \sim_{\text{con}}^a \tilde{\theta}$ and $\theta \sim_{\text{con}}^z \tilde{\theta}$.*

Proof The “only if” part of the statement is trivial. We now show the “if” part.

Let $v := f^{-1} \circ \tilde{f}$. Since $\theta \sim_{\text{con}}^a \tilde{\theta}$, we have that $\tilde{G}^a = PG^a$ and $v = c \circ P^\top$ where P a permutation matrix and c is G^a -preserving. Since $\theta \sim_{\text{con}}^z \tilde{\theta}$, we have that $\tilde{G}^z = \bar{P}G^z \bar{P}^\top$ and $v = \bar{c} \circ \bar{P}^\top$ where $\bar{P}$ is a permutation matrix and $\bar{c}$ is G^z -preserving and $(G^z)^\top$ -preserving. Let C and $\bar{C}$ be the dependency graphs c and $\bar{c}$, respectively.

Choose an arbitrary z . Since $Dc(z)$ is invertible, Lemma 2 implies that there exists a permutation P_0 such that $P_0^\top \subseteq Dv(z)$, which in turn implies that $P_0^\top \subseteq \bar{C}$. Because $P_0^\top \subseteq \bar{C}$, we have that $P_0^\top$ is G^z -preserving and $(G^z)^\top$ -preserving (Proposition 3). By closure under composition and inversion, $\bar{c} \circ P_0$ is G^z - and $(G^z)^\top$ -preserving.

Note that

$$\begin{aligned} c \circ P^\top &= \bar{c} \circ \bar{P}^\top \\ \implies CP^\top &= \bar{C}\bar{P}^\top \\ CP^\top \bar{P} &= \bar{C} \supseteq P_0^\top \\ \implies C &\supseteq P_0^\top \bar{P}^\top P. \end{aligned}$$

This means the permutation $P_0^\top \bar{P}^\top P$ must be G^a -preserving since C is (Proposition 3).

This further implies that $CP^\top \bar{P}P_0 = \bar{C}P_0$ is G^a -preserving by closure under multiplication (recall C is G^a -preserving too). Hence $\bar{c} \circ P_0$ is G^a -preserving.

We thus have that $v = (\bar{c} \circ P_0)(\bar{P}P_0)^\top$ where $(\bar{c} \circ P_0)$ is G^a -, G^z - and $(G^z)^\top$ -preserving. The only thing left to show is that $(\bar{P}P_0)G^a = \tilde{G}^a$ and that $(\bar{P}P_0)G^z(\bar{P}P_0)^\top = \tilde{G}^z$. The former holds since

$$(\bar{P}P_0)G^a = P(P^\top \bar{P}P_0)G^a = PG^a = \tilde{G}^a,$$

where the second equality leverages the fact that $P^\top \bar{P}P_0$ is G^a -preserving. Furthermore,

$$(\bar{P}P_0)G^z(\bar{P}P_0)^\top = \bar{P}G^z(\bar{P}P_0)^\top = \bar{P}G^z P_0^\top \bar{P}^\top = \bar{P}(P_0(G^z)^\top)^\top \bar{P}^\top = \bar{P}G^z \bar{P}^\top = \tilde{G}^z,$$

where the first and fourth equalities leveraged the fact that P_0 is G^z - and $(G^z)^\top$ -preserving. $\blacksquare$

A.5 Technical Lemmas Leading to Theorems 1, 2 & 3

The goal of this section is to introduce and prove Lemma 12 which was crucial in proofs of Theorems 1, 2 & 3. To prove it, we need a few more results, which we present next.

The following two lemmas are standard, but we provide them with proofs for completeness.

Lemma 8 *Let $f : \mathbb{R}^n \rightarrow \mathbb{R}$ be continuous and $A \subseteq \mathbb{R}^n$. If, for all $x \in A$, $f(x) = 0$, then the equality holds on $\overline{A}$.*

Proof We have that $A \subseteq f^{-1}(\{0\})$. Since $\{0\}$ is closed, $f^{-1}(\{0\})$ is also closed by continuity of f . This means $\overline{A} \subseteq f^{-1}(\{0\})$ (since the closure of A is the smallest closed set containing A). ■

Lemma 9 *Let μ be the Lebesgue measure on $\mathbb{R}^{d_z}$ and let $E_0 \subseteq \mathbb{R}^{d_z}$ be a zero measure set, i.e. $\mu(E_0) = 0$. Then, $\overline{\mathbb{R}^{d_z} \setminus E_0} = \mathbb{R}^{d_z}$.*

Proof Clearly, $\overline{\mathbb{R}^{d_z} \setminus E_0} \subseteq \mathbb{R}^{d_z}$.

We now show that $\mathbb{R}^{d_z} \subseteq \overline{\mathbb{R}^{d_z} \setminus E_0}$. Take $z_0 \in \mathbb{R}^{d_z}$ and let U be an open set of $\mathbb{R}^{d_z}$ containing z_0 . Every nonempty open sets have nonzero Lebesgue measure, so

$$0 \neq \mu(U) = \mu(U \cap \mathbb{R}^{d_z}) = \mu(U \cap (\mathbb{R}^{d_z} \setminus E_0)) \implies U \cap (\mathbb{R}^{d_z} \setminus E_0) \neq \emptyset.$$

Since U was arbitrary, this means $z_0 \in \overline{\mathbb{R}^{d_z} \setminus E_0}$. ■

This simple lemma will come in handy when proving Lemma 11.

Lemma 10 *If a permutation P is not G -preserving and C is a G -preserving matrix, we have that $CP^\top$ and $P^\top C$ have a zero on their diagonal.*

Proof Assume P is not G -preserving, hence there exists i, j such that $G_{i,\cdot} \not\subseteq G_{j,\cdot}$, but $P_{i,j} = 1$. Now note that

$$(CP^\top)_{i,i} = C_{i,\cdot}(P_{i,\cdot})^\top = C_{i,\cdot}e_j = C_{i,j},$$

which is equal to zero because C is G -preserving and $G_{i,\cdot} \not\subseteq G_{j,\cdot}$. Similarly,

$$(P^\top C)_{j,j} = (P_{\cdot,j})^\top C_{\cdot,j} = e_i^\top C_{\cdot,j} = C_{i,j} = 0.$$

which concludes the proof. ■

The following lemma is the same as Lemma 12 which is used to prove Theorems 1, 2 & 3, except it does not take into account the “almost everywhere” subtlety. Lemma 12 will extend it to deal with this difficulty.

Lemma 11 *Let $G \in \{0,1\}^{m \times n}$, let $\mathcal{Z}$ be a connected subset of some topological space and let $L : \mathcal{Z} \rightarrow \mathbb{R}^{m \times m}$ be a continuous function such that $L(z)$ is invertible for all $z \in \mathcal{Z}$. Suppose that, for all $z \in \mathcal{Z}$, there exists a permutation matrix $P(z)$ such that $L(z)P(z)$ is G -preserving. Then, there exists a permutation matrix P such that, for all $z \in \mathcal{Z}$, $L(z)P$ is G -preserving.*

Proof The goal of this lemma is to show that in the statement above, one can change the order of the “for all $z \in \mathcal{Z}$ ” and “there exists a permutation”. To do that, we show that if $\mathcal{Z}$ is connected and the map $L(\cdot)$ is continuous, then one can find a single permutation that works for all $z \in \mathcal{Z}$.

Let $\mathcal{G}$ be the set of $\mathbf{G}^a$ -preserving matrices. Recall that, by Proposition 3, $\mathcal{G}$ corresponds to all matrices that have some set of entries equal to zero.

First, since $\mathcal{Z}$ is connected and L is continuous, its image, $L(\mathcal{Z})$, must be connected (by (Munkres, 2000, Theorem 23.5)).

Second, from the hypothesis of the lemma, we know that

$$L(\mathcal{Z}) \subseteq \mathcal{L} := \left(\bigcup_{\pi \in \mathfrak{S}_m} \mathcal{G}P_\pi \right) \setminus \{\text{singular matrices}\},$$

where $\mathfrak{S}_m$ is the set of permutations and $\mathcal{G}P_\pi = \{LP_\pi \mid L \in \mathcal{G}\}$. We can rewrite the set $\mathcal{L}$ above as

$$\mathcal{L} = \left(\bigcup_{\pi \in \mathfrak{S}_m} \mathcal{G}P_\pi \setminus \{\text{singular matrices}\} \right).$$

We now define an equivalence relation $\sim$ over permutations: $\pi \sim \pi'$ iff $P_\pi P_{\pi'}^\top$ is $\mathbf{G}$ -preserving. One can verify that the relation $\sim$ is indeed an equivalence relation by using the fact that invertible $\mathbf{G}$ -preserving matrices form a group (Proposition 4). We notice that

$$\pi \sim \pi' \implies \mathcal{G} = \mathcal{G}P_\pi P_{\pi'}^\top \implies \mathcal{G}P_{\pi'} = \mathcal{G}P_\pi, \quad (53)$$

where the first implication holds because $\mathbf{G}$ -preserving matrices are closed under matrix multiplication (Proposition 4). Let $\mathfrak{S}_m / \sim$ be the set of equivalence classes induced by $\sim$ and let Π stand for one such equivalence class. Thanks to (53), we can define, for all $\Pi \in \mathfrak{S}_m / \sim$, the following set:

$$V_\Pi := \mathcal{G}P_\pi \setminus \{\text{singular matrices}\}, \text{ for some } \pi \in \Pi,$$

where the specific choice of $\pi \in \Pi$ is arbitrary (any $\pi' \in \Pi$ would yield the same definition, by (53)). This construction allows us to write

$$\mathcal{L} = \bigcup_{\Pi \in \mathfrak{S}_m / \sim} V_\Pi.$$

We now show that $\{V_\Pi\}_{\Pi \in \mathfrak{S}_m / \sim}$ forms a partition of $\mathcal{L}$. Choose two distinct equivalence classes of permutations Π and Π' and let $\pi \in \Pi$ and $\pi' \in \Pi'$ be representatives. We will now prove that

$$\mathcal{G}P_\pi \cap \mathcal{G}P_{\pi'} \subseteq \{\text{singular matrices}\}. \quad (54)$$

To achieve this, we proceed by contradiction: Suppose there exists an invertible matrix $A \in \mathcal{G}P_\pi \cap \mathcal{G}P_{\pi'}$. By Lemma 2, there exists a permutation P s.t. $AP^\top$ has no zero on its diagonal. Of course, the permutation P belongs to only one equivalence class and thus either $P \notin \Pi$ or $P \notin \Pi'$. Without loss of generality, assume the former. We thus have that $P \not\sim \pi$ and thus $P_\pi P^\top$ is not $\mathbf{G}$ -preserving. We can thus write

$$A \in \mathcal{G}P_\pi \cap \mathcal{G}P_{\pi'} \implies AP^\top \in \mathcal{G}P_\pi P^\top \cap \mathcal{G}P_{\pi'} P^\top.$$

By Lemma 10, all matrices in $\mathcal{G}P_\pi P^\top$ have a zero on their diagonal. This is a contradiction with $AP^\top \in \mathcal{G}P_\pi P^\top$, since, as we said, $AP^\top$ has no zero on its diagonal. We thus conclude that no invertible matrix is in the intersection $\mathcal{G}P_\pi \cap \mathcal{G}P_{\pi'}$ and thus (54) holds.

We thus have that

$$V_\Pi \cap V_{\Pi'} = \emptyset,$$

which shows that $\{V_\Pi\}_{\Pi \in \mathfrak{S}_m/\sim}$ is indeed a partition of $\mathcal{L}$.

Each V_Π is closed in $\mathcal{L}$ (w.r.t. the subset topology inherited from $\mathbb{R}^{m \times m}$) since

$$V_\Pi = \mathcal{G}P_\pi \setminus \{\text{singular matrices}\} = \mathcal{L} \cap \underbrace{\mathcal{G}P_\pi}_{\text{closed in } \mathbb{R}^{m \times m}}.$$

Moreover, V_Π is open in $\mathcal{L}$, since

$$V_\Pi = \mathcal{L} \setminus \underbrace{\bigcup_{\Pi' \neq \Pi} V_{\Pi'}}_{\text{closed in } \mathcal{L}}.$$

Thus, for any $\Pi \in \mathfrak{S}(\mathcal{B})/\sim$, the sets V_Π and $\bigcup_{\Pi' \neq \Pi} V_{\Pi'}$ forms a *separation* (see (Munkres, 2000, Section 23)). Since $L(\mathcal{Z})$ is a connected subset of $\mathcal{L}$, it must lie completely in V_Π or $\bigcup_{\Pi' \neq \Pi} V_{\Pi'}$, by (Munkres, 2000, Lemma 23.2). Since this is true for all Π , it must follow that there exists a Π^* such that $L(\mathcal{Z}) \subseteq V_{\Pi^*}$. Choose any representative $P_* \in \Pi^*$. We thus have that, for all $z \in \mathcal{Z}$, $L(z) = C(z)P_*^\top$, where $C(z)$ is G -preserving, which completes the proof. ■

The goal of the next result is to relax the conditions of Lemma 11 so that $L(z)P(z)$ is G -preserving for *almost all* $z \in \mathbb{R}^{d_z}$.

Lemma 12 *Let $G \in \{0, 1\}^{m \times n}$ and let $L : \mathbb{R}^{d_z} \rightarrow \mathbb{R}^{m \times m}$ be a continuous function such that $L(z)$ is invertible for all $z \in \mathbb{R}^{d_z}$. Suppose that, for almost all $z \in \mathbb{R}^{d_z}$ (i.e. except on a set E_0 of Lebesgue measure zero), there exists a permutation matrix $P(z)$ such that $L(z)P(z)$ is G -preserving. Then, there exists a permutation matrix P such that, for all $z \in \mathbb{R}^{d_z}$, $L(z)P$ is G -preserving.*

Proof We know that for all $z \in \mathbb{R}^{d_z} \setminus E_0$, where $\mu(E_0) = 0$ (Lebesgue measure zero), there exists a permutation matrix $P(z)$ such that $L(z)P(z)$ is G -preserving. For all permutations P , define $\mathcal{Z}^{(P)} := \{z \in \mathbb{R}^{d_z} \setminus E_0 \mid P(z) = P\}$. The collection of all sets $\mathcal{Z}^{(P)}$ is finite (since there are finitely many permutations) and forms a partition of $\mathbb{R}^{d_z} \setminus E_0$. Of course, for all P , $L(z)P$ is G -preserving for all $z \in \mathcal{Z}^{(P)}$. By Lemma 8, we can extend this statement to the closure, i.e. for all $z \in \overline{\mathcal{Z}^{(P)}}$, $L(z)P$ is G -preserving.

Furthermore, we have that $\bigcup_P \overline{\mathcal{Z}^{(P)}} = \overline{\bigcup_P \mathcal{Z}^{(P)}} = \overline{\mathbb{R}^{d_z} \setminus E_0} = \mathbb{R}^{d_z}$, where the first equality is a standard property of closure (which holds only for finite unions), and the last equality holds by Lemma 9. We thus have that, for all $z \in \mathbb{R}^{d_z}$, there exists a permutation $P(z)$ such that $L(z)P(z)$ is G -preserving. Since $\mathbb{R}^{d_z}$ is connected we can apply Lemma 11 to get the desired conclusion. ■

A.6 Connecting to the Graphical Criterion of [Lachapelle et al. \(2022\)](#)

The goal of this section is to prove Proposition 7 which states if some graphical criterion holds (Assumption 5), then $\theta \sim_{\text{con}}^{z,a} \hat{\theta}$ implies $\theta \sim_{\text{perm}} \hat{\theta}$, i.e. complete disentanglement. We recall Assumption 5.

Assumption 5 (Graphical criterion, [Lachapelle et al. \(2022\)](#)) Let $G = [G^z G^a]$ be a graph. For all $i \in \{1, \dots, d_z\}$,

$$\left(\bigcap_{j \in \text{Ch}_i^z} \text{Pa}_j^z \right) \cap \left(\bigcap_{j \in \text{Pa}_i^z} \text{Ch}_j^z \right) \cap \left(\bigcap_{\ell \in \text{Pa}_i^a} \text{Ch}_\ell^a \right) = \{i\},$$

where Pa_i^z and Ch_i^z are the sets of parents and children of node z_i in G^z , respectively, while Ch_ℓ^a is the set of children of a_ℓ in G^a .

We note that the above assumption is slightly different from the original one from [Lachapelle et al. \(2022\)](#), since the intersections run over Ch_i^z , Pa_i^z and Pa_i^a instead of over some sets of indexes $\mathcal{I}, \mathcal{J} \subseteq \{1, \dots, d_z\}$ and $\mathcal{L} \subseteq \{1, \dots, d_a\}$. This slightly simplified criterion is equivalent to the original one, which we now demonstrate for the interested reader.

Proposition 11 Let $G = [G^z G^a] \in \{0, 1\}^{d_z \times (d_z + d_a)}$. The criterion of Assumption 5 holds for G if and only if the following holds for G : For all $i \in \{1, \dots, d_z\}$, there exist sets $\mathcal{I}, \mathcal{J} \subseteq \{1, \dots, d_z\}$ and $\mathcal{L} \subseteq \{1, \dots, d_a\}$ such that

$$\left(\bigcap_{j \in \mathcal{I}} \text{Pa}_j^z \right) \cap \left(\bigcap_{j \in \mathcal{J}} \text{Ch}_j^z \right) \cap \left(\bigcap_{\ell \in \mathcal{L}} \text{Ch}_\ell^a \right) = \{i\},$$

Proof The direction “ $\implies$ ” is trivial, since we can simply choose $\mathcal{I} := \text{Ch}_i^z$, $\mathcal{J} := \text{Pa}_i^z$ and $\mathcal{L} := \text{Pa}_i^a$.

To show the other direction, we notice that we must have $\mathcal{I} \subseteq \text{Ch}_i^z$, $\mathcal{J} \subseteq \text{Pa}_i^z$ and $\mathcal{L} \subseteq \text{Pa}_i^a$, otherwise one of the sets in the intersection would not contain i , contradicting the criterion. Thus, the criterion of Def. 5 intersects the same sets or more sets. Moreover, these potential additional sets must contain i because of the obvious facts that $j \in \text{Ch}_i^z \iff i \in \text{Pa}_j^z$ and $\ell \in \text{Pa}_i^a \iff i \in \text{Ch}_\ell^a$, thus they do not change the result of the intersection. ■

To prove Proposition 7, we will need the following lemma.

Lemma 13 Let $G \in \{0, 1\}^{m \times n}$ and c be a diffeomorphism with a dependency graph given by $C \in \{0, 1\}^{m \times m}$ (Definition 2). The function c is G -preserving (Definition 12) if and only if

$$\forall i, C_{i,\cdot} \subseteq \bigcap_{k \in G_{i,\cdot}} G_{\cdot,k}.$$

Proof We leverage Proposition 3.

$$\begin{aligned} G_{i,\cdot} \not\subseteq G_{j,\cdot} &\iff \exists k \text{ s.t. } G_{i,k} = 1 \text{ and } G_{j,k} = 0 \\ &\iff \exists k \in G_{i,\cdot} \text{ s.t. } j \notin G_{\cdot,k} \\ &\iff j \notin \bigcap_{k \in G_{i,\cdot}} G_{\cdot,k}. \end{aligned} \tag{55}$$

■

Proposition 7 (Complete disentanglement as a special case) *Let $\theta := (\mathbf{f}, p, \mathbf{G})$ and $\hat{\theta} := (\hat{\mathbf{f}}, \hat{p}, \hat{\mathbf{G}})$ be two models satisfying Assumptions 1, 2 & 3. If $\theta \sim_{\text{con}}^{\mathbf{z}, \mathbf{a}} \hat{\theta}$ and $\mathbf{G}$ satisfies Assumption 5, then $\theta \sim_{\text{perm}} \hat{\theta}$.*

Proof By definition of $\sim_{\text{con}}^{\mathbf{z}, \mathbf{a}}$, we know that the entanglement graph for $(\mathbf{f}, \hat{\mathbf{f}})$ is given by $\mathbf{V} = \mathbf{C}\mathbf{P}^\top$ where $\mathbf{P}$ is a permutation and $\mathbf{C}$ is a binary matrix that is $\mathbf{G}^a$ -preserving, $\mathbf{G}^z$ -preserving and $(\mathbf{G}^z)^\top$ -preserving. Using Lemma 13, we have that, for all i ,

$$\begin{aligned} C_{i,\cdot} &\subseteq \left(\bigcap_{j \in \mathbf{G}_{i,\cdot}^z} \mathbf{G}_{\cdot,j}^z \right) \cap \left(\bigcap_{j \in \mathbf{G}_{\cdot,i}^z} \mathbf{G}_{j,\cdot}^z \right) \cap \left(\bigcap_{j \in \mathbf{G}_{i,\cdot}^a} \mathbf{G}_{\cdot,j}^a \right) \\ &= \left(\bigcap_{j \in \mathbf{Pa}_i^z} \mathbf{Ch}_j^z \right) \cap \left(\bigcap_{j \in \mathbf{Ch}_i^z} \mathbf{Pa}_j^z \right) \cap \left(\bigcap_{\ell \in \mathbf{Pa}_i^a} \mathbf{Ch}_\ell^a \right) \\ &= \{i\}. \end{aligned}$$

Thus, $\mathbf{C}$ is in fact the identity matrix, and hence $\theta \sim_{\text{perm}} \hat{\theta}$. ■

Appendix B. Identifiability Theory - Exponential Family Case

B.1 Technical Lemmas and Definitions

We recall the definition of a minimal sufficient statistic in an exponential family, which can be found in [Wainwright and Jordan \(2008, p. 40\)](#).

Definition 20 (Minimal sufficient statistic) *Given a parameterized distribution in the exponential family, as in (30), we say that its sufficient statistic $\mathbf{s}_i$ is minimal when there is no $v \neq 0$ such that $v^\top \mathbf{s}_i(z)$ is constant for all $z \in \mathcal{Z}$.*

The following Lemma gives a characterization of minimality, which will be useful in the proof of Theorem 4.

Lemma 14 (Characterization of minimal $\mathbf{s}$) *A sufficient statistic of an exponential family distribution $\mathbf{s} : \mathcal{Z} \rightarrow \mathbb{R}^k$ is minimal if and only if there exists $\mathbf{z}_{(0)}, \mathbf{z}_{(1)}, \dots, \mathbf{z}_{(k)}$ belonging to the support $\mathcal{Z}$ such that the following k -dimensional vectors are linearly independent:*

$$\mathbf{s}(\mathbf{z}_{(1)}) - \mathbf{s}(\mathbf{z}_{(0)}), \dots, \mathbf{s}(\mathbf{z}_{(k)}) - \mathbf{s}(\mathbf{z}_{(0)}). \quad (56)$$

Proof. We start by showing the “if” part of the statement. Suppose there exist $\mathbf{z}_{(0)}, \dots, \mathbf{z}_{(k)}$ in $\mathcal{Z}$ such that the vectors of (56) are linearly independent. By contradiction, suppose that $\mathbf{s}$ is not minimal, i.e. there exist a nonzero vector v and a scalar b such that $v^\top \mathbf{s}(z) = b$ for all $z \in \mathcal{Z}$. Notice that

$b = v^\top s(z_{(0)})$. Hence, $v^\top (s(z_{(i)}) - s(z_{(0)})) = 0$ for all $i = 1, \dots, k$. This can be rewritten in matrix form as

$$v^\top [s(z_{(1)}) - s(z_{(0)}) \dots s(z_{(k)}) - s(z_{(0)})] = 0,$$

which implies that the matrix in the above equation is not invertible. This is a contradiction.

We now show the “only if” part of the statement. Suppose that there is no $z_{(0)}, \dots, z_{(k)}$ such that the vectors of (56) are linearly independent. Choose an arbitrary $z_{(0)} \in \mathcal{Z}$. We thus have that $U := \text{span}\{s(z) - s(z_{(0)}) \mid z \in \mathcal{Z}\}$ is a proper subspace of $\mathbb{R}^k$. This means the orthogonal complement of U , $U^\perp$, has dimension 1 or greater. We can thus pick a nonzero vector $v \in U^\perp$ such that $v^\top (s(z) - s(z_{(0)})) = 0$ for all $z \in \mathcal{Z}$, which is to say that $v^\top s(z)$ is constant for all $z \in \mathcal{Z}$, and thus, s is not minimal. ■

B.2 Proof of Linear Identifiability (Theorem 4)

Theorem 4 (Conditions for linear identifiability - Adapted from Khemakhem et al. (2020a)) Let $\theta := (f, \lambda, G)$ and $\hat{\theta} := (\hat{f}, \hat{\lambda}, \hat{G})$ be two models satisfying Assumptions 1, 2 & 9. Further assume that

1. **[Observational equivalence]** $\theta \sim_{\text{obs}} \hat{\theta}$ (Definition 4);
2. **[Minimal sufficient statistics]** For all i , the sufficient statistic s_i is minimal (see below).
3. **[Sufficient variability]** The natural parameter λ varies “sufficiently” as formalized by Assumption 10 (see below).

Then, $\theta \sim_{\text{lin}} \hat{\theta}$ (Def. 17).

Proof First, we apply Proposition 2 to get

$$\tilde{p}(z^t \mid z^{<t}, a^{<t}) = p(v(z^t) \mid v(z^{<t}), a^{<t}) |\det Dv(z^t)|. \quad (57)$$

Linear relationship between $s(f^{-1}(x))$ and $s(\hat{f}^{-1}(x))$. By taking the logarithm on each side of (57) and making explicit the exponential family form, we get

$$\begin{aligned} & \sum_{i=1}^{d_z} \log h_i(z_i^t) + s_i(z_i^t)^\top \lambda_i(G_i^z \odot z^{<t}, G_i^a \odot a^{<t}) - \psi_i(z^{<t}, a^{<t}) \\ &= \sum_{i=1}^{d_z} \log h_i(v_i(z^t)) + s_i(v_i(z^t))^\top \hat{\lambda}_i(\hat{G}_i^z \odot v(z^{<t}), \hat{G}_i^a \odot a^{<t}) - \hat{\psi}_i(v(z^{<t}), a^{<t}) \\ & \quad + \log |\det Dv(z^t)| \end{aligned} \quad (58)$$

Note that (58) holds for all $z^{<t}$ and $a^{<t}$. In particular, we evaluate it at the points given in the assumption of sufficient variability of Theorem 4. We evaluate the equation at $(z^t, z_{(r)}, a_{(r)})$ and

$(\mathbf{z}^t, \mathbf{z}_{(0)}, \mathbf{a}_{(0)})$ and take the difference which yields¹⁰

$$\begin{aligned} & \sum_{i=1}^{d_z} \mathbf{s}_i(\mathbf{z}_i^t)^\top [\boldsymbol{\lambda}_i(\mathbf{G}_i^z \odot \mathbf{z}_{(r)}, \mathbf{G}_i^a \odot \mathbf{a}_{(r)}) - \boldsymbol{\lambda}_i(\mathbf{G}_i^z \odot \mathbf{z}_{(0)}, \mathbf{G}_i^a \odot \mathbf{a}_{(0)})] - \psi_i(\mathbf{z}_{(r)}, \mathbf{a}_{(r)}) + \psi_i(\mathbf{z}_{(0)}, \mathbf{a}_{(0)}) \\ &= \sum_{i=1}^{d_z} \mathbf{s}_i(\mathbf{v}_i(\mathbf{z}^t))^\top [\hat{\boldsymbol{\lambda}}_i(\hat{\mathbf{G}}_i^z \odot \mathbf{v}(\mathbf{z}_{(r)}), \hat{\mathbf{G}}_i^a \odot \mathbf{a}_{(r)}) - \hat{\boldsymbol{\lambda}}_i(\hat{\mathbf{G}}_i^z \odot \mathbf{v}(\mathbf{z}_{(0)}), \hat{\mathbf{G}}_i^a \odot \mathbf{a}_{(0)})] \\ & \quad - \hat{\psi}_i(\mathbf{v}(\mathbf{z}_{(r)}), \mathbf{a}_{(r)}) + \hat{\psi}_i(\mathbf{v}(\mathbf{z}_{(0)}), \mathbf{a}_{(0)}) \end{aligned} \quad (59)$$

We regroup all normalization constants ψ into a term $d(\mathbf{z}_{(r)}, \mathbf{z}_{(0)}, \mathbf{a}_{(r)}, \mathbf{a}_{(0)})$ and write

$$\begin{aligned} & \mathbf{s}(\mathbf{z}^t)^\top [\boldsymbol{\lambda}(\mathbf{z}_{(r)}, \mathbf{a}_{(r)}) - \boldsymbol{\lambda}(\mathbf{z}_{(0)}, \mathbf{a}_{(0)})] \\ &= \mathbf{s}(\mathbf{v}(\mathbf{z}^t))^\top [\hat{\boldsymbol{\lambda}}(\mathbf{v}(\mathbf{z}_{(r)}), \mathbf{a}_{(r)}) - \hat{\boldsymbol{\lambda}}(\mathbf{v}(\mathbf{z}_{(0)}), \mathbf{a}_{(0)})] + d(\mathbf{z}_{(r)}, \mathbf{z}_{(0)}, \mathbf{a}_{(r)}, \mathbf{a}_{(0)}) . \end{aligned} \quad (60)$$

Define

$$\begin{aligned} \mathbf{w}_{(r)} &:= \boldsymbol{\lambda}(\mathbf{z}_{(r)}, \mathbf{a}_{(r)}) - \boldsymbol{\lambda}(\mathbf{z}_{(0)}, \mathbf{a}_{(0)}) \\ \hat{\mathbf{w}}_{(r)} &:= \hat{\boldsymbol{\lambda}}(\mathbf{v}(\mathbf{z}_{(r)}), \mathbf{a}_{(r)}) - \hat{\boldsymbol{\lambda}}(\mathbf{v}(\mathbf{z}_{(0)}), \mathbf{a}_{(0)}) \\ d_{(r)} &:= d(\mathbf{z}_{(r)}, \mathbf{z}_{(0)}, \mathbf{a}_{(r)}, \mathbf{a}_{(0)}) , \end{aligned}$$

which yields

$$\mathbf{s}(\mathbf{z}^t)^\top \mathbf{w}_{(r)} = \mathbf{s}(\mathbf{v}(\mathbf{z}^t))^\top \hat{\mathbf{w}}_{(r)} + d_{(r)} . \quad (61)$$

We can regroup the $\mathbf{w}_{(r)}$ into a matrix and the $d_{(r)}$ into a vector:

$$\begin{aligned} \mathbf{W} &:= [\mathbf{w}_{(1)} \dots \mathbf{w}_{(kd_z)}] \in \mathbb{R}^{kd_z \times kd_z} \\ \hat{\mathbf{W}} &:= [\hat{\mathbf{w}}_{(1)} \dots \hat{\mathbf{w}}_{(kd_z)}] \in \mathbb{R}^{kd_z \times kd_z} \\ \mathbf{d} &:= [d_{(1)} \dots d_{(kd_z)}] \in \mathbb{R}^{1 \times kd_z} . \end{aligned}$$

Since (61) holds for all $1 \leq p \leq kd_z$, we can write

$$\mathbf{s}(\mathbf{z}^t)^\top \mathbf{W} = \mathbf{s}(\mathbf{v}(\mathbf{z}^t))^\top \hat{\mathbf{W}} + \mathbf{d} .$$

Note that $\mathbf{W}$ is invertible by the assumption of variability, hence

$$\mathbf{s}(\mathbf{z}^t)^\top = \mathbf{s}(\mathbf{v}(\mathbf{z}^t))^\top \hat{\mathbf{W}} \mathbf{W}^{-1} + \mathbf{d} \mathbf{W}^{-1} .$$

Let $\mathbf{b} := (\mathbf{d} \mathbf{W}^{-1})^\top$ and $\mathbf{L} := (\hat{\mathbf{W}} \mathbf{W}^{-1})^\top$. We can thus rewrite this as

$$\mathbf{s}(\mathbf{z}^t) = \mathbf{L} \mathbf{s}(\mathbf{v}(\mathbf{z}^t)) + \mathbf{b} . \quad (62)$$

10. Note that $\mathbf{z}_{(0)}$ and $\mathbf{z}_{(r)}$ can have different dimensionalities if they come from different time steps. It is not an issue to combine equations from different time steps, since (58) holds for all values of t , $\mathbf{z}^t$, $\mathbf{z}^{<t}$ and $\mathbf{a}^{<t}$.

Invertibility of L . We now show that L is invertible. By Lemma 14, the fact that the s_i are minimal is equivalent to, for all $i \in \{1, \dots, d_z\}$, having elements $z_i^{(0)}, \dots, z_i^{(k)}$ in $\mathcal{Z}$ such that the family of vectors

$$s_i(z_i^{(1)}) - s_i(z_i^{(0)}), \dots, s_i(z_i^{(k)}) - s_i(z_i^{(0)})$$

is linearly independent. Define

$$z^{(0)} := [z_1^{(0)} \dots z_{d_z}^{(0)}]^\top \in \mathbb{R}^{d_z}$$

For all $i \in \{1, \dots, d_z\}$ and all $p \in \{1, \dots, k\}$, define the vectors

$$z^{(p,i)} := [z_1^{(0)} \dots z_{i-1}^{(0)} z_i^{(p)} z_{i+1}^{(0)} \dots z_{d_z}^{(0)}]^\top \in \mathbb{R}^{d_z}.$$

For a specific $1 \leq p \leq k$ and $i \in \{1, \dots, d_z\}$, we can take the following difference based on (62)

$$s(z^{(p,i)}) - s(z^{(0)}) = L[s(v(z^{(p,i)})) - s(v(z^{(0)}))], \quad (63)$$

where the left hand side is a vector filled with zeros except for the block corresponding to $s_i(z_i^{(p,i)}) - s_i(z_i^{(0)})$. Let us define

$$\begin{aligned} \Delta s^{(i)} &:= [s(z^{(1,i)}) - s(z^{(0)}) \dots s(z^{(k,i)}) - s(z^{(0)})] \in \mathbb{R}^{kd_z \times k} \\ \Delta \hat{s}^{(i)} &:= [s(v(z^{(1,i)})) - s(v(z^{(0)})) \dots s(v(z^{(k,i)})) - s(v(z^{(0)}))] \in \mathbb{R}^{kd_z \times k}. \end{aligned}$$

Note that the columns of $\Delta s^{(i)}$ are linearly independent and all rows are filled with zeros except for the block of rows $\{(i-1)k+1, \dots, ik\}$. We can thus rewrite (63) in matrix form

$$\Delta s^{(i)} = L \Delta \hat{s}^{(i)}.$$

We can regroup these equations for every i by doing

$$[\Delta s^{(1)} \dots \Delta s^{(d_z)}] = L[\Delta \hat{s}^{(1)} \dots \Delta \hat{s}^{(d_z)}]. \quad (64)$$

Notice that the newly formed matrix on the left hand side has size $kd_z \times kd_z$ and is block diagonal. Since every block is invertible, the left hand side of (64) is an invertible matrix, which in turn implies that L is invertible. This completes the proof. $\blacksquare$

B.3 Proof of Theorem 5

Lemma 15 Let $\theta := (f, p, G)$ satisfy Assumptions 1, 2, 3, 4 & 9 and let $q := \log p$. Then

$$\begin{aligned} D_z^t q(z^t \mid z^{<t}, a^{<t}) &= \lambda(z^{<t}, a^{<t})^\top Ds(z^t) + D(\log h)(z^t) \\ H_{z,a}^{t,\tau} q(z^t \mid z^{<t}, a^{<t}) &= Ds(z^t)^\top D_a^\tau \lambda(z^{<t}, a^{<t}) \\ H_{z,z}^{t,\tau} q(z^t \mid z^{<t}, a^{<t}) &= Ds(z^t)^\top D_z^\tau \lambda(z^{<t}, a^{<t}). \end{aligned}$$

Proof We have

$$\begin{aligned} \log p(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t}) &:= \log h(\mathbf{z}^t) + \mathbf{s}(\mathbf{z}^t)^\top \boldsymbol{\lambda}(\mathbf{z}^{<t}, \mathbf{a}^{<t}) - \psi(\mathbf{z}^{<t}, \mathbf{a}^{<t}) \\ &\quad \log h(\mathbf{z}^t) + \boldsymbol{\lambda}(\mathbf{z}^{<t}, \mathbf{a}^{<t})^\top \mathbf{s}(\mathbf{z}^t) - \psi(\mathbf{z}^{<t}, \mathbf{a}^{<t}). \end{aligned}$$

We can differentiate the above w.r.t. $\mathbf{z}^t$ to get

$$D_z^t q(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t}) = \boldsymbol{\lambda}(\mathbf{z}^{<t}, \mathbf{a}^{<t})^\top D\mathbf{s}(\mathbf{z}^t) + D(\log h)(\mathbf{z}^t)$$

Differentiating the above w.r.t. $\mathbf{z}^\tau$ or $\mathbf{a}^\tau$ yields the desired result. $\blacksquare$

Theorem 5 (Disentanglement via sparse temporal dependencies in exponential families) *Let $\boldsymbol{\theta} := (\mathbf{f}, \boldsymbol{\lambda}, \mathbf{G})$ and $\hat{\boldsymbol{\theta}} := (\hat{\mathbf{f}}, \hat{\boldsymbol{\lambda}}, \hat{\mathbf{G}})$ be two models satisfying Assumptions 1, 2, 3, 4, 9 as well as all assumptions of Theorem 4. Further suppose that*

1. *The sufficient statistic $\mathbf{s}$ is d_z -dimensional ($k = 1$) and is a diffeomorphism from $\mathbb{R}^{d_z}$ to $\mathbf{s}(\mathbb{R}^{d_z})$;*
2. *[Sufficient influence of $\mathbf{z}$] The Jacobian of the ground-truth transition function $\boldsymbol{\lambda}$ with respect to $\mathbf{z}$ varies “sufficiently”, as formalized in Assumption 11;*

Then, there exists a permutation matrix $\mathbf{P}$ such that $\mathbf{P}\mathbf{G}^z\mathbf{P}^\top \subseteq \hat{\mathbf{G}}^z$. Further assume that

3. *[Sparsity regularization] $\|\hat{\mathbf{G}}^z\|_0 \leq \|\mathbf{G}^z\|_0$;*

Then, $\boldsymbol{\theta} \sim_{\text{con}}^z \hat{\boldsymbol{\theta}}$ (Def. 14) & $\boldsymbol{\theta} \sim_{\text{lin}} \hat{\boldsymbol{\theta}}$ (Def. 17), which together implies that

$$\mathbf{v}(\mathbf{z}) = \mathbf{s}^{-1}(\mathbf{C}\mathbf{P}^\top \mathbf{s}(\mathbf{z}) + \mathbf{b}),$$

where $\mathbf{b} \in \mathbb{R}^{d_z}$ and $\mathbf{C} \in \mathbb{R}^{d_z \times d_z}$ is invertible, $\mathbf{G}^z$ - and $(\mathbf{G}^z)^\top$ -preserving (Definition 11).

Proof Recall the equation we derived in Section 3.2:

$$H_{z,z}^{t,\tau} \hat{q}(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t}) = D\mathbf{v}(\mathbf{z}^t)^\top H_{z,z}^{t,\tau} q(\mathbf{v}(\mathbf{z}^t) \mid \mathbf{v}(\mathbf{z}^{<t}), \mathbf{a}^{<t}) D\mathbf{v}(\mathbf{z}^\tau).$$

Using Lemma 15, we get

$$D\mathbf{s}(\mathbf{z}^t)^\top D_z^\tau \hat{\boldsymbol{\lambda}}(\mathbf{z}^{<t}, \mathbf{a}^{<t}) = D\mathbf{v}(\mathbf{z}^t)^\top D\mathbf{s}(\mathbf{v}(\mathbf{z}^t))^\top D_z^\tau \boldsymbol{\lambda}(\mathbf{v}(\mathbf{z}^{<t}), \mathbf{a}^{<t}) D\mathbf{v}(\mathbf{z}^\tau). \quad (65)$$

Note that Assumption 3 requires that $D_z^\tau \boldsymbol{\lambda}(\mathbf{z}^{<t}, \mathbf{a}^{<t}) \subseteq \mathbf{G}^z$ and $D_z^\tau \hat{\boldsymbol{\lambda}}(\mathbf{z}^{<t}, \mathbf{a}^{<t}) \subseteq \hat{\mathbf{G}}^z$. Theorem 4 implies that there exist an invertible matrix $\mathbf{L} \in \mathbb{R}^{d_z \times d_z}$ and a vector $\mathbf{b} \in \mathbb{R}^{d_z}$ such that

$$\mathbf{v}(\mathbf{z}) = \mathbf{s}^{-1}(\mathbf{L}\mathbf{s}(\mathbf{z}) + \mathbf{b}).$$

Taking the derivative of the above w.r.t. $\mathbf{z}$, we obtain

$$D\mathbf{v}(\mathbf{z}) = D\mathbf{s}^{-1}(\mathbf{L}\mathbf{s}(\mathbf{z}) + \mathbf{b})\mathbf{L}D\mathbf{s}(\mathbf{z}) \quad (66)$$

$$= D\mathbf{s}^{-1}(\mathbf{s}(\mathbf{v}(\mathbf{z})))\mathbf{L}D\mathbf{s}(\mathbf{z}) \quad (67)$$

$$= D\mathbf{s}(\mathbf{v}(\mathbf{z}))^{-1}\mathbf{L}D\mathbf{s}(\mathbf{z}), \quad (68)$$

where we used $\mathbf{s}(\mathbf{v}(\mathbf{z})) = \mathbf{L}\mathbf{s}(\mathbf{z}) + \mathbf{b}$ to go from the first to the second line and used the inverse function theorem to go from the second to the third line. Plugging (68) into (65) yields

$$\begin{aligned} & D\mathbf{s}(\mathbf{z}^t)^\top D_z^\tau \hat{\boldsymbol{\lambda}}(\mathbf{z}^{<t}, \mathbf{a}^{<t}) \\ &= D\mathbf{s}(\mathbf{z}^t)^\top \mathbf{L}^\top D\mathbf{s}(\mathbf{v}(\mathbf{z}^t))^{-\top} D\mathbf{s}(\mathbf{v}(\mathbf{z}^t))^\top D_z^\tau \boldsymbol{\lambda}(\mathbf{v}(\mathbf{z}^{<t}), \mathbf{a}^{<t}) D\mathbf{s}(\mathbf{v}(\mathbf{z}^\tau))^{-1} \mathbf{L} D\mathbf{s}(\mathbf{z}^\tau) \\ &= D\mathbf{s}(\mathbf{z}^t)^\top \mathbf{L}^\top D_z^\tau \boldsymbol{\lambda}(\mathbf{v}(\mathbf{z}^{<t}), \mathbf{a}^{<t}) D\mathbf{s}(\mathbf{v}(\mathbf{z}^\tau))^{-1} \mathbf{L} D\mathbf{s}(\mathbf{z}^\tau), \end{aligned}$$

which implies

$$D\mathbf{s}(\mathbf{z}^t)^\top D_z^\tau \hat{\boldsymbol{\lambda}}(\mathbf{z}^{<t}, \mathbf{a}^{<t}) = D\mathbf{s}(\mathbf{z}^t)^\top \mathbf{L}^\top D_z^\tau \boldsymbol{\lambda}(\mathbf{v}(\mathbf{z}^{<t}), \mathbf{a}^{<t}) D\mathbf{s}(\mathbf{v}(\mathbf{z}^\tau))^{-1} \mathbf{L} D\mathbf{s}(\mathbf{z}^\tau) \quad (69)$$

$$D_z^\tau \hat{\boldsymbol{\lambda}}(\mathbf{z}^{<t}, \mathbf{a}^{<t}) D\mathbf{s}(\mathbf{z}^\tau)^{-1} = \mathbf{L}^\top D_z^\tau \boldsymbol{\lambda}(\mathbf{v}(\mathbf{z}^{<t}), \mathbf{a}^{<t}) D\mathbf{s}(\mathbf{v}(\mathbf{z}^\tau))^{-1} \mathbf{L}, \quad (70)$$

where we right- and left-multiplied by $D\mathbf{s}(\mathbf{z}^t)^{-\top}$ and $D\mathbf{s}(\mathbf{z}^\tau)$, respectively. Let us define

$$\Lambda(\gamma) := D_z^\tau \boldsymbol{\lambda}(\mathbf{v}(\mathbf{z}^{<t}), \mathbf{a}^{<t}) D\mathbf{s}(\mathbf{v}(\mathbf{z}^\tau))^{-1} \quad \hat{\Lambda}(\gamma) := D_z^\tau \hat{\boldsymbol{\lambda}}(\mathbf{v}(\mathbf{z}^{<t}), \mathbf{a}^{<t}) D\mathbf{s}(\mathbf{v}(\mathbf{z}^\tau))^{-1},$$

where $\gamma = (t, \tau, \mathbf{z}^{<t}, \mathbf{a}^{<t})$. Note that because $D\mathbf{s}$ is diagonal, we have that $\Lambda(\gamma) \subseteq \mathbf{G}^z$ and $\hat{\Lambda}(\gamma) \subseteq \hat{\mathbf{G}}^z$. Using this notation, we can rewrite (70) as

$$\underbrace{\hat{\Lambda}(\gamma)}_{\subseteq \hat{\mathbf{G}}^z} = \mathbf{L}^\top \underbrace{\Lambda(\gamma)}_{\subseteq \mathbf{G}^z} \mathbf{L}.$$

Thanks to Assumption 11, we can apply the same argument as in Theorem 5 to show that $\mathbf{L} = \mathbf{C}\mathbf{P}^\top$ where $\mathbf{C}$ is a matrix that is both $\mathbf{G}^z$ -preserving and $(\mathbf{G}^z)^\top$ -preserving, as desired. $\blacksquare$

B.4 Connecting to Sufficient Influence Assumptions of Lachapelle et al. (2022)

In this section, we relate the nonparametric sufficient influence assumptions of this work, i.e. Assumptions 7 & 8, to the analogous assumptions of Lachapelle et al. (2022) for exponential families, i.e. Assumptions 11 & 12, the latter of which we recall below.

Assumption 12 (Sufficient influence of $\mathbf{a}$ (Lachapelle et al., 2022)) Assume $k = 1$, i.e. the sufficient statistics $\mathbf{s}_i$ are one-dimensional. For all $\ell \in \{1, \dots, d_a\}$, there exist $\{(\mathbf{z}_{(r)}, \mathbf{a}_{(r)}, \epsilon_{(r)}, \tau_{(r)})\}_{r=1}^{|\mathbf{Ch}_\ell^a|}$ belonging to their respective support such that

$$\text{span} \left\{ \boldsymbol{\lambda}(\mathbf{z}_{(r)}, \mathbf{a}_{(r)} + \epsilon_{(r)} \mathbf{E}^{(\ell, \tau)}) - \boldsymbol{\lambda}(\mathbf{z}_{(r)}, \mathbf{a}_{(r)}) \right\}_{r=1}^{|\mathbf{Ch}_\ell^a|} = \mathbb{R}_{\mathbf{Ch}_\ell^a}^{d_z},$$

where $\epsilon \in \mathbb{R}$ and $\mathbf{E}^{(\ell, \tau)} \in \mathbb{R}^{d_a \times t}$ is the one-hot matrix with the entry (ℓ, τ) set to one.

The following proposition shows that, when the exponential family holds with $k = 1$, we have that (i) for the “sufficient influence of $\mathbf{a}$ ” assumptions, the nonparametric and exponential family versions are actually equivalent, and (ii) for the “sufficient influence of $\mathbf{z}$ ” assumptions, the nonparametric version implies the exponential family version.

Proposition 12 (Sufficient influence assumptions: nonparametric v.s. exponential) *Let the parameter $\theta := (\mathbf{f}, p, \mathbf{G})$ satisfy Assumptions 1, 2, 3 & 9. Further, assume that $k = 1$ and that $D\mathbf{s}(\mathbf{z}) \in \mathbb{R}^{d_z \times d_z}$ is invertible everywhere. Then,*

Sufficient influence of $\mathbf{a}$: Assumption 7 (nonparametric) $\iff$ Assumption 12 (exponential family)

Sufficient influence of $\mathbf{z}$: Assumption 8 (nonparametric) $\implies$ Assumption 11 (exponential family)

Proof We start by proving the first equivalence for the sufficient influence of $\mathbf{a}$ assumptions. By using Lemma 15 we see that

$$\begin{aligned} & \text{span} \left\{ D_z^{t(r)} \log p(\mathbf{z} \mid \mathbf{z}_{(r)}, \mathbf{a}_{(r)} + \epsilon_{(r)} \mathbf{E}^{(\ell, \tau(r))}) - D_z^{t(r)} \log p(\mathbf{z} \mid \mathbf{z}_{(r)}, \mathbf{a}_{(r)}) \right\}_{r=1}^{|\text{Ch}_\ell^{\mathbf{a}}|} \\ &= \text{span} \left\{ D\mathbf{s}(\mathbf{z})^\top \boldsymbol{\lambda}(\mathbf{z}_{(r)}, \mathbf{a}_{(r)} + \epsilon_{(r)} \mathbf{E}^{(\ell, \tau(r))}) - D\mathbf{s}(\mathbf{z})^\top \boldsymbol{\lambda}(\mathbf{z}_{(r)}, \mathbf{a}_{(r)}) \right\}_{r=1}^{|\text{Ch}_\ell^{\mathbf{a}}|} \\ &= D\mathbf{s}(\mathbf{z})^\top \text{span} \left\{ \boldsymbol{\lambda}(\mathbf{z}_{(r)}, \mathbf{a}_{(r)} + \epsilon_{(r)} \mathbf{E}^{(\ell, \tau(r))}) - \boldsymbol{\lambda}(\mathbf{z}_{(r)}, \mathbf{a}_{(r)}) \right\}_{r=1}^{|\text{Ch}_\ell^{\mathbf{a}}|}. \end{aligned} \quad (71)$$

We start by showing “ $\Leftarrow$ ”. Assumption 12 implies that (71) is equal to $D\mathbf{s}(\mathbf{z}^t)^\top \mathbb{R}_{\text{Ch}_\ell^{\mathbf{a}}}^{d_z}$ which is equal to $\mathbb{R}_{\text{Ch}_\ell^{\mathbf{a}}}^{d_z}$ since $D\mathbf{s}(\mathbf{z}^t)$ is invertible everywhere and is diagonal. To show “ $\Rightarrow$ ”, we can apply the same argument.

We now show that Assumption 8 implies Assumption 11. we again use Lemma 15 and see that

$$\begin{aligned} \mathbb{R}_{\mathbf{G}^z}^{d_z} &= \text{span} \left\{ H_{z,z}^{t(r), \tau(r)} \log p(\mathbf{z} \mid \mathbf{z}_{(r)}, \mathbf{a}_{(r)}) \right\}_{r=1}^{||\mathbf{G}^z||_0} \\ &= \text{span} \left\{ D\mathbf{s}(\mathbf{z})^\top D_z^{\tau(r)} \boldsymbol{\lambda}(\mathbf{z}_{(r)}, \mathbf{a}_{(r)}) \right\}_{r=1}^{||\mathbf{G}^z||_0} \\ &= D\mathbf{s}(\mathbf{z})^\top \text{span} \left\{ D_z^{\tau(r)} \boldsymbol{\lambda}(\mathbf{z}_{(r)}, \mathbf{a}_{(r)}) D\mathbf{s}(\mathbf{z})^{-1} \right\}_{r=1}^{||\mathbf{G}^z||_0} D\mathbf{s}(\mathbf{z}). \end{aligned}$$

Now recall that, in Assumption 8, we had that $\mathbf{z} = \mathbf{z}^{\tau(r)}$ for all $r = 1, \dots, ||\mathbf{G}^z||_0$, which allows us to write

$$\mathbb{R}_{\mathbf{G}^z}^{d_z} = D\mathbf{s}(\mathbf{z})^\top \text{span} \left\{ D_z^{\tau(r)} \boldsymbol{\lambda}(\mathbf{z}_{(r)}, \mathbf{a}_{(r)}) D\mathbf{s}(\mathbf{z}^{\tau(r)})^{-1} \right\}_{r=1}^{||\mathbf{G}^z||_0} D\mathbf{s}(\mathbf{z}),$$

which implies

$$\text{span} \left\{ D_z^{\tau(r)} \boldsymbol{\lambda}(\mathbf{z}_{(r)}, \mathbf{a}_{(r)}) D\mathbf{s}(\mathbf{z}^{\tau(r)})^{-1} \right\}_{r=1}^{||\mathbf{G}^z||_0} = D\mathbf{s}(\mathbf{z})^{-\top} \mathbb{R}_{\mathbf{G}^z}^{d_z} D\mathbf{s}(\mathbf{z})^{-1} = \mathbb{R}_{\mathbf{G}^z}^{d_z},$$

where the last equality holds because $D\mathbf{s}(\mathbf{z})$ is diagonal and invertible everywhere. ■

Appendix C. Experiments

C.1 Synthetic Datasets

We now provide a detailed description of the synthetic datasets used in the experiments of Section 8.

For all experiments, the dimensionality of $\mathbf{x}^t$ is $d_x = 20$ and the ground-truth $\mathbf{f}$ is a random neural network with three hidden layers of 20 units with Leaky-ReLU activations with a negative slope of 0.2. The weight matrices are sampled according to a 0-1 Gaussian distribution and, to ensure that $\mathbf{f}$ is injective, as assumed in all theorems of this paper, we orthogonalize its columns. Inspired by typical weight initialization in NN (Glorot and Bengio, 2010), we rescale the weight matrices by $\sqrt{\frac{2}{1+0.2^2}} \sqrt{\frac{2}{d_{in}+d_{out}}}$. The standard deviation of the Gaussian noise added to $\mathbf{f}(\mathbf{z}^t)$ is set to $\sigma = 10^{-2}$ throughout. Since the goal of the experiments is to validate our identifiability results, which assume infinite data, all datasets considered here are very large: 1 million examples.

We now present the different choices of ground-truth $p(\mathbf{z}^t | \mathbf{z}^{<t}, \mathbf{a}^{<t})$ we explored in our experiments. In all cases considered (except for the experiment with $k = 2$ of Table 3), it is a Gaussian with covariance $0.0001I$ independent of $(\mathbf{z}^{<t}, \mathbf{a}^{<t})$ and a mean given by some function $\mu(\mathbf{z}^{t-1}, \mathbf{a}^{t-1})$. Notice that we hence are in the case where $k = 1$ with monotonic sufficient statistics, which is not covered by the theory of Khemakhem et al. (2020a). Throughout, we set $d_z = 10$ and, unless explicitly specified otherwise, we set $d_a = 10$. In all *Time* datasets, sequences have length $T = 2$. In *Action* datasets, the value of T has no consequence since we assume there is no time dependence.

C.1.1 DATASETS SATISFYING THE GRAPHICAL CRITERION

The datasets of this section satisfy the graphical criterion of Section 3.7. This means our theory predicts complete disentanglement (Definition 7). Unless specified otherwise, all datasets satisfy their respective sufficient influence assumptions (Section 3.8). These can be checked using Remark 5 combined with standard facts about independence of the sine and cosine functions.

ActionDiag (Figure 5). In this dataset, $d_a = d_x$ and the connectivity matrix between $\mathbf{a}^{t-1}$ and $\mathbf{z}^t$ is diagonal, which trivially implies that the graphical criterion of Section 3.7 is satisfied. The mean function is given by

$$\mu(\mathbf{z}^{t-1}, \mathbf{a}^{t-1}) := \sin(\mathbf{a}^{t-1}),$$

where $\sin$ is applied element-wise. Moreover, the components of the action vector $\mathbf{a}^{t-1}$ are sampled independently and uniformly between -2 and 2 . The same sampling scheme is used for all following datasets. One can check that the sufficient influence assumption (Assumption 6) holds.

ActionNonDiag (Figure 5). We consider a case where the graphical criterion of Section 3.7 is satisfied non-trivially. Let

$$\mathbf{G}^a := \begin{pmatrix} 1 & & & 1 \\ 1 & 1 & & \\ & 1 & \ddots & \\ & & \ddots & 1 \\ & & & 1 & 1 \end{pmatrix} \quad (72)$$

be the adjacency matrix between $\mathbf{a}^{t-1}$ and $\mathbf{z}^t$. The i th row, denoted by $\mathbf{G}_i^a$, corresponds to the parents of $\mathbf{z}_i^t$ in $\mathbf{a}^{t-1}$. Note that it is analogous to the graph shown in Figure 4, which satisfies the

graphical criterion. The mean function is given by

$$\mu(\mathbf{z}^{t-1}, \mathbf{a}^{t-1}) := \begin{bmatrix} \mathbf{G}_1^{\mathbf{a}} \cdot \sin(\frac{3}{\pi} \mathbf{a}^{t-1}) \\ \mathbf{G}_2^{\mathbf{a}} \cdot \sin(\frac{4}{\pi} \mathbf{a}^{t-1} + 1) \\ \vdots \\ \mathbf{G}_{d_z}^{\mathbf{a}} \cdot \sin(\frac{d_z+2}{\pi} \mathbf{a}^{t-1} + d_z - 1) \end{bmatrix}. \quad (73)$$

One can check that the sufficient influence assumption (Assumption 6) holds, due to the independence of sines with different frequencies.

ActionNonDiag_{NoSuffInf} (Table 3). This dataset has the same ground truth adjacency matrix as the above dataset (72), but a different transition function which does not satisfy the assumption of sufficient influence (Section 6). We sampled a matrix $\mathbf{W}$ with independent Normal 0-1 entries. The mean function is thus

$$\mu(\mathbf{z}^{t-1}, \mathbf{a}^{t-1}) := (\mathbf{G}^{\mathbf{a}} \odot \mathbf{W}) \mathbf{a}^{t-1},$$

where $\odot$ is the Hadamard product (a.k.a. element-wise product).

ActionNonDiag_{k=2} (Table 3). This dataset has the “double diagonal” adjacency matrix of (72) and the same mean function of (73), but the variance of $\mathbf{z}^t$ (we assume diagonal covariance) depends on $\mathbf{a}^{t-1}$ via

$$\sigma^2(\mathbf{z}^{t-1}, \mathbf{a}^{t-1}) := \frac{1}{10d_a} \begin{bmatrix} \exp(\mathbf{G}_1^{\mathbf{a}} \cdot \cos(\frac{3}{\pi} \mathbf{a}^{t-1})) \\ \exp(\mathbf{G}_2^{\mathbf{a}} \cdot \cos(\frac{4}{\pi} \mathbf{a}^{t-1} + 1)) \\ \vdots \\ \exp(\mathbf{G}_{d_z}^{\mathbf{a}} \cdot \cos(\frac{d_z+2}{\pi} \mathbf{a}^{t-1} + d_z - 1)) \end{bmatrix}. \quad (74)$$

TimeDiag (Figure 5). In this dataset, each $\mathbf{z}_i^t$ has only $\mathbf{z}_i^{t-1}$ as parent. This trivially satisfies the graphical criterion of Section 3.7. The mean function is given by

$$\mu(\mathbf{z}^{t-1}, \mathbf{a}^{t-1}) := \mathbf{z}^{t-1} + 0.5 \sin(\mathbf{z}^{t-1}),$$

where the sin function is applied element-wise. Notice that no auxiliary variables are required. One can check that the sufficient variability assumption (Assumption 10) and sufficient influence assumption (Assumption 11) of Theorem 5 (exponential family) holds.

TimeNonDiag (Figure 5). We consider a case where the graphical criterion of Section 3.7 is satisfied non-trivially. Let

$$\mathbf{G}^z := \begin{pmatrix} 1 & & & & \\ 1 & 1 & & & \\ \vdots & & \ddots & & \\ 1 & & & 1 & \\ 1 & 1 & \dots & 1 & 1 \end{pmatrix} \quad (75)$$

be the adjacency matrix between $\mathbf{z}^t$ and $\mathbf{z}^{t-1}$. The i th row of $\mathbf{G}^z$, denoted by $\mathbf{G}_i^z$, corresponds to the parents of $\mathbf{z}_i^t$. Notice that this connectivity matrix has no 2-cycles and all self-loops are present.

Thus, by Proposition 8, it satisfies the graphical criterion of Section 3.7. The mean function in this case is given by

$$\mu(\mathbf{z}^{t-1}, \mathbf{a}^{t-1}) := \mathbf{z}^{t-1} + 0.5 \begin{bmatrix} \mathbf{G}_1^z \cdot \sin(\frac{3}{\pi} \mathbf{z}^{t-1}) \\ \mathbf{G}_2^z \cdot \sin(\frac{4}{\pi} \mathbf{z}^{t-1} + 1) \\ \vdots \\ \mathbf{G}_{d_z}^z \cdot \sin(\frac{d_z+2}{\pi} \mathbf{z}^{t-1} + d_z - 1) \end{bmatrix}, \quad (76)$$

which is analogous to (73). One can verify that this transition model satisfies the sufficient variability assumption (Assumption 10) and sufficient influence assumption (Assumption 10) of Theorem 5 (exponential family) holds.

TimeNonDiag_{NoSuffInf} (Table 3). This dataset has the same ground truth adjacency matrix as in (75), but a different transition function that does not satisfy the assumption of sufficient influence. We sampled a transition matrix W with independent Normal 0-1 entries. The transition function is thus

$$\mu(\mathbf{z}^{t-1}, \mathbf{a}^{t-1}) := \mathbf{z}^{t-1} + 0.5(\mathbf{G}^z \odot W) \mathbf{z}^{t-1}.$$

TimeNonDiag_{k=2} (Table 3). This dataset has the lower triangular adjacency matrix of (75) and the same mean function of (76), but the variance of $\mathbf{z}^t$ (we assume diagonal covariance) depends on $\mathbf{z}^{t-1}$ via

$$\sigma^2(\mathbf{z}^{t-1}, \mathbf{a}^{t-1}) := \frac{1}{10d_z} \begin{bmatrix} \exp(\mathbf{G}_1^z \cdot \cos(\frac{3}{\pi} \mathbf{z}^{t-1})) \\ \exp(\mathbf{G}_2^z \cdot \cos(\frac{4}{\pi} \mathbf{z}^{t-1} + 1)) \\ \vdots \\ \exp(\mathbf{G}_{d_z}^z \cdot \cos(\frac{d_z+2}{\pi} \mathbf{z}^{t-1} + d_z - 1)) \end{bmatrix}. \quad (77)$$

C.1.2 DATASETS THAT VIOLATE THE GRAPHICAL CRITERION

The transition mechanisms for the action and temporal datasets of Figure 7 and Table 4 are (73) and (76), respectively, except for the graphs which are different.

ActionBlockDiag and ActionBlockNonDiag (Figure 7). The left graph corresponds to ActionBlockDiag while the right one corresponds to ActionBlockNonDiag.

$$\mathbf{G}_{(1)}^a := \begin{bmatrix} 1 & & & & & \\ 1 & & & & & \\ & 1 & & & & \\ & 1 & & & & \\ & & 1 & & & \\ & & 1 & & & \\ & & & 1 & & \\ & & & 1 & & \\ & & & & 1 & \\ & & & & & 1 \end{bmatrix} \quad \mathbf{G}_{(2)}^a := \begin{bmatrix} 1 & & & & & \\ 1 & & & & & \\ & 1 & & & & \\ & 1 & & & & \\ & & 1 & & & \\ & & 1 & & & \\ & & & 1 & & \\ & & & 1 & & \\ & & & & 1 & \\ & 1 & & & & 1 \\ 1 & & & & & 1 \end{bmatrix}$$

TimeBlockDiag and TimeBlockNonDiag (Figure 7). The left graph corresponds to TimeBlockDiag while the right one corresponds to TimeBlockNonDiag.

$$G_{(1)}^z := \begin{bmatrix} 1 & 1 & & & & & & & & \\ 1 & 1 & & & & & & & & \\ & & 1 & 1 & & & & & & \\ & & 1 & 1 & & & & & & \\ & & & & 1 & 1 & & & & \\ & & & & 1 & 1 & & & & \\ & & & & & & 1 & 1 & & \\ & & & & & & 1 & 1 & & \\ & & & & & & & & 1 & 1 \\ & & & & & & & & 1 & 1 \end{bmatrix} \quad G_{(2)}^z := \begin{bmatrix} 1 & 1 & & & & & & & 1 & 1 \\ 1 & 1 & & & & & & & 1 & 1 \\ & & 1 & 1 & & & & & & \\ & & 1 & 1 & & & & & & \\ & & & & 1 & 1 & & & & \\ & & & & 1 & 1 & & & & \\ & & & & & & 1 & 1 & & \\ & & & & & & 1 & 1 & & \\ & & & & & & & & 1 & 1 \\ 1 & 1 & & & 1 & 1 & & & 1 & 1 \\ 1 & 1 & & & 1 & 1 & & & 1 & 1 \end{bmatrix}$$

ActionRandomGraphs and TimeRandomGraphs (Table 4). The transition mechanisms are the same as in the ActionNonDiag and TimeNonDiag datasets, i.e. they are given by (73) & (76), respectively. However, the graphs are sampled randomly, with various levels of sparsity. For the ActionRandomGraphs dataset, we have $G_{i,j}^a \sim \text{Ber}(p)$ and independent. For the TimeRandomGraphs datasets, it is the same except for the diagonal elements, which are forced to be active, i.e. $G_{i,i}^z = 1$.

C.2 Implementation Details of our Constrained VAE Method

All details of our implementation match those of Lachapelle et al. (2022) (except for the constrained optimization introduced in Section 5).

Learned mechanisms. Every coordinate z_i of the latent vector has its own mechanism $\hat{p}(z_i^t | z^{<t}, a^{<t})$ that is Gaussian with mean outputted by $\hat{\mu}_i(z^{t-1}, a^{t-1})$ (a multilayer perceptron with 5 layers of 512 units) and a learned variance which does not depend on the previous time steps. For learning, we use the typical parameterization of the Gaussian distribution with μ and σ^2 and not its exponential family parameterization. Throughout, the dimensionality of z^t in the learned model always matches the dimensionality of the ground-truth (same for the baselines). Learning the dimensionality of z^t is left for future work.

Prior of z^1 in time-sparsity experiments. In *time-sparsity* experiments, the prior of the first latent $\hat{p}(z^1)$ (when $t = 1$) is modeled separately as a Gaussian with learned mean and learned diagonal covariance. Note that this learned covariance at time $t = 1$ is different from the subsequent learned conditional covariance at time $t > 1$.

Encoder/Decoder. In all experiments, including baselines, both the encoder and the decoder is modeled by a neural network with 6 fully connected hidden layers of 512 units with LeakyReLU activation with negative slope 0.2. For all VAE-based methods, the encoder outputs the mean and a diagonal covariance. Moreover, $p(x|z)$ has a *learned* isotropic covariance $\sigma^2 I$. Note that $\sigma^2 I$ corresponds to the covariance of the independent noise n^t in the equation $x^t = f(z^t) + n^t$.

C.3 Baselines

In synthetic experiments of Sec. 8, all methods used a minibatch size of 1024 and the same encoder and decoder architecture: A MLP with 6 layers of 512 units with LeakyReLU activations (negative slope of 0.2). We manually tuned the learning rate of each method to ensure proper convergence. For

VAE-based methods, i.e. TCVAE, SlowVAE and iVAE, we are always choosing $p(x|z)$ Gaussian with a covariance $\sigma^2 I$ and learn σ^2 .

β -TCVAE. We used the implementation provided in the original paper by [Chen et al. \(2018\)](#), which is available at <https://github.com/rtqichen/beta-tcvae>. We used a learning rate of $1e-4$.

iVAE. We used the implementation available at <https://github.com/ilkhem/icebeam> from [Khemakhem et al. \(2020a\)](#). In it, the mean of the prior $p(z|a)$ is fixed to zero, while its diagonal covariance is allowed to depend on a through an MLP. We change this to allow the mean to also depend on a through the neural network (with 5 layers and width 512). We also lower bounded its variance as well as the variance of $q(z | x, a)$ to improve the stability of learning. In the original implementation, the covariance of $p(x|z)$ was not learned. We found that learning it (analogously to what we do in our method) improved performance. We used a learning rate of $1e-4$.

SlowVAE. We used the implementation provided in https://github.com/bethgelab/slow_disentanglement ([Klindt et al., 2021](#)). Like for other VAE-based methods, we modeled $p(x|z)$ as a Gaussian with covariance $\sigma^2 I$ and learned σ^2 .

PCL. We used the implementation provided here: https://github.com/bethgelab/slow_disentanglement/tree/baselines. PCL ([Hyvarinen and Morioka, 2017](#)) stands for “permutation contrastive learning” and works as follows: Given sequential data $\{x^t\}_{t=1}^T$, PCL trains a regression function $r((x', x))$ to discriminate between pairs of adjacent observations (positive pairs) and randomly matched pairs (negative pairs). The regression function has the form

$$r((x, x')) = \sum_{i=1}^{d_z} B_i(h_i(x), h_i(x')),$$

where $h : \mathbb{R}^{d_x} \rightarrow \mathbb{R}^{d_z}$ is the encoder and $B_i : \mathbb{R}^2 \rightarrow \mathbb{R}$ are learned functions. In our implementation, the functions B_i are fully connected neural networks with 5 layers and 512 hidden units. We experimented with the less expressive function suggested in the original work, but found that the extra capacity improved performance across all datasets we considered.

C.4 Unsupervised Hyperparameter Selection

In practice, MCC cannot be measured since the ground-truth latent variables are not observed. Unlike in standard machine learning settings, hyperparameter selection for disentanglement cannot be performed simply by evaluating goodness of fit on a validation set and selecting the highest scoring model, since there is usually a trade-off between goodness of fit and disentanglement ([Locatello et al., 2019](#), Sec. 5.4). To avoid this problem, [Duan et al. \(2020\)](#) introduced *an unsupervised disentanglement ranking* (UDR), which, for all combinations of hyperparameters, measures how consistent are different random initializations of the algorithm. The authors argue that hyperparameters yielding disentangled representation typically yields consistent representations. In our experiments, the consistency of a given hyperparameter combination is measured as follows: for every pair of models, we compute the MCC between their representations. Then, we report the median of all pairwise MCC. This gives a UDR score for every hyperparameter values considered. Figure 9 reports the ELBO (normalized between zero and one), the MCC and the UDR score for the experiments of Figure 5. We can visualize the trade-off between ELBO and MCC. That being said, MCC and UDR correlate nicely except for the TimeNonDiag dataset, in which this correlation breaks for stronger regularization. We noticed that these specific runs correspond to excessively

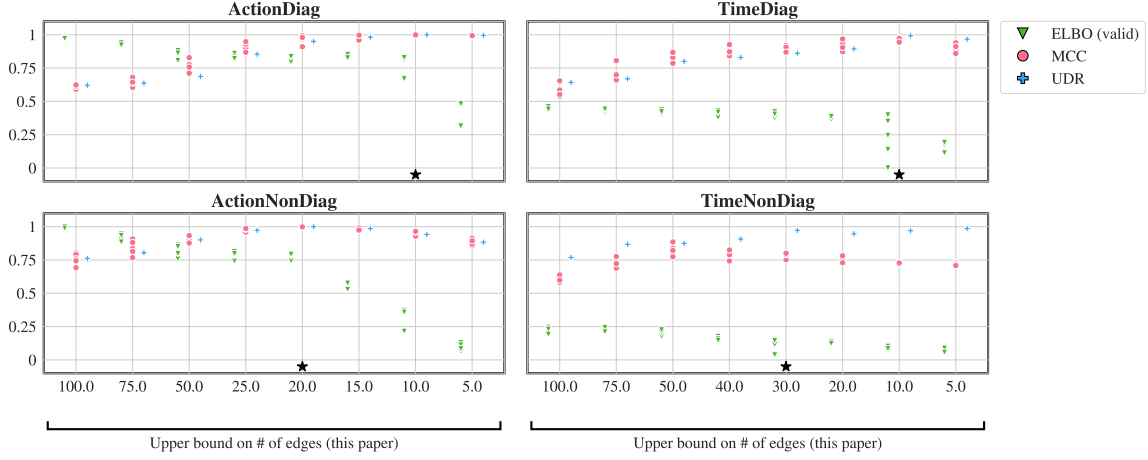


Figure 9: Investigating the link between goodness of fit (ELBO), disentanglement (MCC) and UDR. The ELBO is normalized so that it remains between 0 and 1.

sparse graph, with fewer than 10 edges (out of 100 possible edges). The black star indicates the hyperparameter selected by UDR when excluding coefficient values which yields graphs with less than 10 edges (on average).

Baselines. Two of the baselines considered had hyperparameters to tune, SlowVAE (Klindt et al., 2021) and TCVAE (Chen et al., 2018). For SlowVAE, we did a grid search on the following values, $\gamma \in \{1.0, 2.0, 4.0, 8.0, 16.0\}$ and $\alpha \in \{1, 3, 6, 10\}$. For TCVAE, we explored $\beta \in \{1, 2, 3, 4, 5\}$ but the optimal value in terms of disentanglement was almost always 1. Values of β larger than 5 led to instabilities during training. The hyperparameters were selected using UDR, as described in the paragraph above.

Appendix D. Miscellaneous

D.1 On the Invertibility of the Mixing Function

Throughout this work as well as many others (Hyvarinen and Morioka, 2016, 2017; Hyvärinen et al., 2019; Khemakhem et al., 2020a; Locatello et al., 2020; Klindt et al., 2021), it is assumed that the mixing function mapping the latent factors to the observation is a diffeomorphism onto its image. In this section, we briefly discuss the practical implications of this assumption.

Recall that a diffeomorphism is a differentiable bijective function with a differentiable inverse. We start by addressing the bijective part of the assumption. To understand it, we consider a plausible situation where the mapping f is not invertible. Consider the minimal example of Fig. 1 consisting of a tree, a robot and a ball. Assume that the ball can be hidden behind either the tree or the robot. Then, the mixing function f is not invertible because, given only the image, it is impossible to know whether the ball is behind the tree or the robot. Thus, this situation is not covered by our theory. Intuitively, one could infer, at least approximately, where the ball is hidden based on previous time frames. Allowing for this form of occlusion is left as future work. See also Mansouri et al. (2022) for further discussion about how one can relax this assumption.

We believe the differentiable part of this assumption is only a technicality that could probably be relaxed to being piecewise differentiable. Our experiments were performed with data generated

with a piecewise linear f , which is not differentiable only on a set of (Lebesgue) measure zero, but this was not an issue in practice.

D.2 Contrasting with the Assumptions of Khemakhem et al. (2020a) & Yao et al. (2022b)

In this section, we discuss two identifiability results previously proposed in the literature that do not leverage sparsity (Khemakhem et al., 2020a; Yao et al., 2022a). We show that these results do not apply to the simple homoscedastic Gaussian latent models of the form $p(\mathbf{z}_i^t | \mathbf{z}^{t-1}) = \mathcal{N}(\mathbf{z}_i^t | \boldsymbol{\mu}_i(\mathbf{z}^{t-1}), \sigma_i^2)$, contrary to our theory, as we saw in Examples 9, 10 and 12. We will see that in the context of a Gaussian latent model, both results require the variance to vary sufficiently strongly. We believe that such a requirement is not well suited for nearly deterministic environments such as the one depicted in Figure 1.

Khemakhem et al. (2020a). The most significant distinction between the theory of Khemakhem et al. (2020a) (iVAE) and ours is how identifiability up to permutation is obtained: Theorems 2 & 3 from iVAE shows that if the assumptions of their Theorem 1 (which is essentially Theorem 4) are satisfied and $\mathbf{s}_i$ has dimension $k > 1$ or is non-monotonic, then the model is not just identifiable up to linear transformation but up to permutations (and rescalings). In contrast, our theory covers the case where $k = 1$ and $\mathbf{s}_i$ is monotonic, like in the homoscedastic Gaussian case. Interestingly, Khemakhem et al. (2020a) mentioned this specific case as a counterexample to their theory in their Proposition 3. The extra power of our theory comes from the extra *structure* in the dependencies of the latent factors coupled with sparsity regularization. We note that, assuming the latent factors are Gaussian, the variability assumption of Theorem 4 combined with $k > 1$ requires the variance to vary sufficiently, which is implausible in the nearly deterministic environment of Figure 1.

Yao et al. (2022a). This work (Theorem 1) requires that, for each value of $\mathbf{z}^t$, the $2d_z$ functions

$$\frac{\partial^2}{\partial \mathbf{z}_i^t \partial \mathbf{z}^{t-1}} \log p(\mathbf{z}_i^t | \mathbf{z}^{t-1}) \text{ and } \frac{\partial^3}{(\partial \mathbf{z}_i^t)^2 \partial \mathbf{z}^{t-1}} \log p(\mathbf{z}_i^t | \mathbf{z}^{t-1}) \text{ for } i = 1 \dots d_z,$$

seen as functions from $\mathbb{R}^{d_z}$ to $\mathbb{R}^{d_z}$ are linearly independent. Indeed, if $p(\mathbf{z}_i^t | \mathbf{z}^{t-1}) = \mathcal{N}(\mathbf{z}_i^t | \boldsymbol{\mu}_i(\mathbf{z}^{t-1}), \sigma_i^2)$, one can easily derive that

$$\begin{aligned} \frac{\partial}{\partial \mathbf{z}_i^t} \log p(\mathbf{z}_i^t | \mathbf{z}_i^{t-1}) &= -(\mathbf{z}_i - \boldsymbol{\mu}_i(\mathbf{z}^{t-1}))/\sigma_i^2 \\ \frac{\partial^2}{(\partial \mathbf{z}_i^t)^2} \log p(\mathbf{z}_i^t | \mathbf{z}_i^{t-1}) &= -1/\sigma_i^2 \\ \frac{\partial^2}{(\partial \mathbf{z}_i^t)^2 \partial \mathbf{z}^{t-1}} \log p(\mathbf{z}_i^t | \mathbf{z}_i^{t-1}) &= \mathbf{0}, \end{aligned}$$

which shows that the assumption of Yao et al. (2022a, Theorem 1) does not hold for homoscedastic Gaussian latent models. We further notice that, had the variance σ_i^2 depend on $\mathbf{z}^{t-1}$, the identifiability result of Yao et al. (2022b) could have applied.

D.3 Derivation of the ELBO

In this section, we derive the evidence lower bound presented in Section 5.

$$\begin{aligned} \log p(\mathbf{x}^{\leq T} \mid \mathbf{a}^{<T}) &= \\ \mathbb{E}_{q(\mathbf{z}^{\leq T} \mid \mathbf{x}^{\leq T}, \mathbf{a}^{<T})} &\left[\log \frac{q(\mathbf{z}^{\leq T} \mid \mathbf{x}^{\leq T}, \mathbf{a}^{<T})}{p(\mathbf{z}^{\leq T} \mid \mathbf{x}^{\leq T}, \mathbf{a}^{<T})} \right. \\ &\quad \left. + \log \frac{p(\mathbf{z}^{\leq T}, \mathbf{x}^{\leq T} \mid \mathbf{a}^{<T})}{q(\mathbf{z}^{\leq T} \mid \mathbf{x}^{\leq T}, \mathbf{a}^{<T})} \right] \end{aligned} \quad (78)$$

$$\begin{aligned} &\geq \mathbb{E}_{q(\mathbf{z}^{\leq T} \mid \mathbf{x}^{\leq T}, \mathbf{a}^{<T})} \left[\log \frac{p(\mathbf{z}^{\leq T}, \mathbf{x}^{\leq T} \mid \mathbf{a}^{<T})}{q(\mathbf{z}^{\leq T} \mid \mathbf{x}^{\leq T}, \mathbf{a}^{<T})} \right] \\ &= \mathbb{E}_{q(\mathbf{z}^{\leq T} \mid \mathbf{x}^{\leq T}, \mathbf{a}^{<T})} [\log p(\mathbf{x}^{\leq T} \mid \mathbf{z}^{\leq T}, \mathbf{a}^{<T})] \end{aligned} \quad (79)$$

$$- KL(q(\mathbf{z}^{\leq T} \mid \mathbf{x}^{\leq T}, \mathbf{a}^{<T}) \parallel p(\mathbf{z}^{\leq T} \mid \mathbf{a}^{<T})) \quad (80)$$

where the inequality holds because the term at (78) is a Kullback-Leibler divergence, which is greater or equal to 0. Notice that

$$p(\mathbf{x}^{\leq T} \mid \mathbf{z}^{\leq T}, \mathbf{a}^{<T}) = p(\mathbf{x}^{\leq T} \mid \mathbf{z}^{\leq T}) = \prod_{t=1}^T p(\mathbf{x}^t \mid \mathbf{z}^t). \quad (81)$$

Recall that we are considering a variational posterior of the following form:

$$q(\mathbf{z}^{\leq T} \mid \mathbf{x}^{\leq T}, \mathbf{a}^{<T}) := \prod_{t=1}^T q(\mathbf{z}^t \mid \mathbf{x}^t). \quad (82)$$

Equations (81) & (82) allow us to rewrite the term in (79) as

$$\sum_{t=1}^T \mathbb{E}_{\mathbf{z}^t \sim q(\cdot \mid \mathbf{x}^t)} [\log p(\mathbf{x}^t \mid \mathbf{z}^t)]$$

Notice further that

$$p(\mathbf{z}^{\leq T} \mid \mathbf{a}^{<T}) = \prod_{t=1}^T p(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t}). \quad (83)$$

Using (82) & (83), the KL term (80) can be broken down as a sum of KL as:

$$\sum_{t=1}^T \mathbb{E}_{\mathbf{z}^{<t} \sim q(\cdot \mid \mathbf{x}^{<t})} KL(q(\mathbf{z}^t \mid \mathbf{x}^t) \parallel p(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t}))$$

Putting all together yields the desired ELBO:

$$\begin{aligned} \log p(\mathbf{x}^{\leq T} \mid \mathbf{a}^{<T}) &\geq \sum_{t=1}^T \mathbb{E}_{\mathbf{z}^t \sim q(\cdot \mid \mathbf{x}^t)} [\log p(\mathbf{x}^t \mid \mathbf{z}^t)] \\ &\quad - \mathbb{E}_{\mathbf{z}^{<t} \sim q(\cdot \mid \mathbf{x}^{<t})} KL(q(\mathbf{z}^t \mid \mathbf{x}^t) \parallel p(\mathbf{z}^t \mid \mathbf{z}^{<t}, \mathbf{a}^{<t})). \end{aligned} \quad (84)$$

Appendix E. Author Contributions

E.1 Contributions to the Extended Version

Sébastien Lachapelle developed the idea, the theory and proofs behind mechanism sparsity regularization for disentanglement, wrote the crux of the paper, developed the regularized VAE-based method, and performed most of the experiments. **Rémi Le Priol** provided valuable feedback on the clarity of the manuscript. **Simon Lacoste-Julien** helped with the overall presentation of the paper, clarified the conceptual framework and the motivation, and provided supervision.

E.2 Contributions to the CLear Version (Lachapelle et al., 2022)

Sébastien Lachapelle developed the idea, the theory and proofs behind mechanism sparsity regularization for disentanglement, wrote the first draft of the paper, and designed and implemented the regularized VAE-based method. **Pau Rodríguez López** ran all experiments appearing in the paper, produced associated figures, and ran experiments with image data that are still work in progress. **Yash Sharma** contributed to the research process, the experimental design in particular, implemented and ran experiments on image data that did not make it in the final version, and contributed to the writing and the literature review. **Katie Everett** implemented and ran experiments on image data that did not make it in the final version and contributed to the writing and figures. **Rémi Le Priol** reviewed the proofs of main theorems, simplified some arguments and the general presentation of the proof, and contributed to the writing and figures. **Alexandre Lacoste** produced image datasets that did not make it into the final version and provided supervision. **Simon Lacoste-Julien** helped with the overall presentation of the paper, clarified the conceptual framework and the motivation, and provided supervision.

References

- K. Ahuja, J. Hartford, and Y. Bengio. Properties from mechanisms: an equivariance perspective on identifiable representation learning. In *International Conference on Learning Representations*, 2022a.
- K. Ahuja, J. Hartford, and Y. Bengio. Weakly supervised representation learning with sparse perturbations. *arXiv preprint arXiv:2206.01101*, 2022b.
- K. Ahuja, D. Mahajan, V. Syrgkanis, and I. Mitliagkas. Towards efficient representation identification in supervised learning. In *Conference on Causal Learning and Reasoning*, 2022c.
- K. Ahuja, D. Mahajan, Y. Wang, and Y. Bengio. Interventional causal representation learning. In *International Conference on Machine Learning*, 2023.
- Elias Bareinboim, Juan Correa, Duligur Ibeling, and Thomas Icard. On Pearl’s hierarchy and the foundations of causal inference (1st edition). In *Probabilistic and Causal Inference: the Works of Judea Pearl*. ACM Books, 2022.
- Y. Bengio. The consciousness prior. *arXiv preprint arXiv:1709.08568*, 2019.
- Y. Bengio, A. Courville, and P. Vincent. Representation learning: A review and new perspectives. *IEEE transactions on pattern analysis and machine intelligence*, 2013.

- Y. Bengio, T. Deleu, N. Rahaman, N. R. Ke, S. Lachapelle, O. Bilaniuk, A. Goyal, and C. Pal. A meta-transfer objective for learning to disentangle causal mechanisms. In *International Conference on Learning Representations*, 2020.
- M. Bereket and T. Karaletsos. Modelling cellular perturbations with the sparse additive mechanism shift variational autoencoder. In *Advances in Neural Information Processing Systems*, 2023.
- S. Bing, U. Ninad, J. Wahl, and J. Runge. Identifying linearly-mixed causal representations from multi-node interventions. *arXiv preprint arXiv:2311.02695*, 2023.
- J. Brady, R. S. Zimmermann, Y. Sharma, B. Schölkopf, J. von Kügelgen, and W. Brendel. Provably learning object-centric representations. In *International Conference on Machine Learning*, 2023.
- J. Brehmer, P. de Haan, P. Lippe, and T. Cohen. Weakly supervised causal representation learning. In *Advances in Neural Information Processing Systems*, 2022.
- P. Brouillard, S. Lachapelle, A. Lacoste, S. Lacoste-Julien, and A. Drouin. Differentiable causal discovery from interventional data. In *Advances in Neural Information Processing Systems*, 2020.
- S. Buchholz, M. Besserve, and B. Schölkopf. Function classes for identifiable nonlinear independent component analysis. In *Advances in Neural Information Processing Systems*, 2022.
- S. Buchholz, G. Rajendran, E. Rosenfeld, B. Aragam, B. Schölkopf, and P. K. Ravikumar. Learning linear causal representations from interventions under general nonlinear mixing. In *Advances Neural Information Processing Systems*, 2023.
- R. T. Q. Chen, X. Li, R. G., and D. Duvenaud. Isolating sources of disentanglement in vaes. In *Advances in Neural Information Processing Systems*, 2018.
- D. M. Chickering. Optimal structure identification with greedy search. In *Journal of Machine Learning Research*, 2003.
- P. Comon. Independent component analysis, a new concept? *Signal Processing*, 1994.
- I. Daunhawer, A. Bizeul, E. Palumbo, A. Marx, and J. E Vogt. Identifiability results for multimodal contrastive learning. In *International Conference on Learning Representations*, 2023.
- S. Duan, L. Matthey, A. Saraiva, N. Watters, C. Burgess, A. Lerchner, and I. Higgins. Unsupervised model selection for variational disentangled representation learning. In *International Conference on Learning Representations*, 2020.
- D. Eaton and K. Murphy. Exact bayesian structure learning from uncertain interventions. In *International Conference on Artificial Intelligence and Statistics*, 2007.
- M. Fumero, F. Wenzel, L. Zancato, A. Achille, E. Rodolà, S. Soatto, B. Schölkopf, and F. Locatello. Leveraging sparse and shared feature activations for disentangled representation learning. In *Advances in Neural Information Processing Systems*, 2023.
- J. Gallego-Posada and J. Ramirez. Cooper: a toolkit for lagrangian-based constrained optimization. <https://github.com/cooper-org/cooper>, 2022.

- J. Gallego-Posada, J. Ramirez De Los Rios, and A. Erraqabi. Flexible learning of sparse neural networks via constrained L_0 regularization. In *NeurIPS Workshop LatinX in AI*, 2021.
- L. Girin, S. Leglaive, X. Bie, J. Diard, T. Hueber, and X. Alameda-Pineda. Dynamical variational autoencoders: A comprehensive review. *arXiv preprint arXiv:2008.12595*, 2020.
- X. Glorot and Y. Bengio. Understanding the difficulty of training deep feedforward neural networks. In *International Conference on Artificial Intelligence and Statistics*, 2010.
- A. Goyal and Y. Bengio. Inductive biases for deep learning of higher-level cognition. *arXiv preprint arXiv:2011.15091*, 2021.
- A. Goyal, A. R. Didolkar, N. R. Ke, C. Blundell, P. Beaudoin, N. Heess, M. C. Mozer, and Y. Bengio. Neural production systems. In *Advances in Neural Information Processing Systems*, 2021a.
- A. Goyal, A. Lamb, J. Hoffmann, S. Sodhani, S. Levine, Y. Bengio, and B. Schölkopf. Recurrent independent mechanisms. In *International Conference on Learning Representations*, 2021b.
- L. Gresele, P. K. Rubenstein, A. Mehrjou, F. Locatello, and B. Schölkopf. The incomplete rosetta stone problem: Identifiability results for multi-view nonlinear ica. In *Conference on Uncertainty in Artificial Intelligence Conference*, 2020.
- L. Gresele, J. Von Kügelgen, V. Stimper, B. Schölkopf, and M. Besserve. Independent mechanism analysis, a new concept? In *Advances in Neural Information Processing Systems*, 2021.
- M. U. Gutmann and A. Hyvärinen. Noise-contrastive estimation of unnormalized statistical models, with applications to natural image statistics. *Journal of Machine Learning Research*, 2012.
- H. Hälvä and A. Hyvärinen. Hidden markov nonlinear ica: Unsupervised learning from nonstationary time series. In *Conference on Uncertainty in Artificial Intelligence (UAI)*, 2020.
- H. Hälvä, S. Le Corff, L. Lehericy, J. So, Y. Zhu, E. Gassiat, and A. Hyvarinen. Disentangling identifiable features from noisy data with structured nonlinear ICA. In *Advances in Neural Information Processing Systems*, 2021.
- A. Hauser and P. Bühlmann. Characterization and greedy learning of interventional markov equivalence classes of directed acyclic graphs. *Journal of Machine Learning Research*, 2012.
- I. Higgins, L. Matthey, A. Pal, C. P. Burgess, X. Glorot, M. Botvinick, S. Mohamed, and A. Lerchner. beta-vae: Learning basic visual concepts with a constrained variational framework. In *International Conference on Learning Representations*, 2017.
- D. Horan, E. Richardson, and Y. Weiss. When is unsupervised disentanglement possible? In *Advances in Neural Information Processing Systems*, 2021.
- A. Hyvarinen and H. Morioka. Unsupervised feature extraction by time-contrastive learning and nonlinear ica. In *Advances in Neural Information Processing Systems*, 2016.
- A. Hyvarinen and H. Morioka. Nonlinear ICA of Temporally Dependent Stationary Sources. In *International Conference on Artificial Intelligence and Statistics*, 2017.

- A. Hyvärinen and P. Pajunen. Nonlinear independent component analysis: Existence and uniqueness results. *Neural Networks*, 1999.
- A. Hyvärinen, J. Karhunen, and E. Oja. *Independent Component Analysis*. Wiley, 2001.
- A. Hyvärinen, H. Sasaki, and R. E. Turner. Nonlinear ica using auxiliary variables and generalized contrastive learning. In *International Conference on Artificial Intelligence and Statistics*, 2019.
- A. Hyvärinen, I. Khemakhem, and H. Morioka. Nonlinear independent component analysis for principled disentanglement in unsupervised deep learning. *arXiv preprint arXiv:2303.16535*, 2023.
- A. Jaber, M. Kocaoglu, K. Shanmugam, and E. Bareinboim. Causal discovery from soft interventions with unknown targets: Characterization and learning. In *Advances in Neural Information Processing Systems*, 2020.
- E. Jang, S. Gu, and B. Poole. Categorical reparameterization with gumbel-softmax. *International Conference on Machine Learning*, 2017.
- Y. Jiang and B. Aragam. Learning nonparametric latent causal graphs with unknown interventions. In *Advances in Neural Information Processing Systems*, 2023.
- C. Jutten and J. Herault. Blind separation of sources, part 1: An adaptive algorithm based on neuromimetic architecture. *Signal Process.*, 1991.
- T. Karaletsos, S. Belongie, and G. Rätsch. Bayesian representation learning with oracle constraints. In *International Conference on Learning Representations*, 2016.
- N. R. Ke, O. Bilaniuk, A. Goyal, S. Bauer, H. Larochelle, C. Pal, and Y. Bengio. Learning neural causal models from unknown interventions. *arXiv preprint arXiv:1910.01075*, 2019.
- N. R. Ke, A. R. Didolkar, S. Mittal, A. Goyal, G. Lajoie, S. Bauer, D. J. Rezende, M. C. Mozer, Y. Bengio, and C. Pal. Systematic evaluation of causal discovery in visual model based reinforcement learning. *arXiv preprint arXiv:2107.00848*, 2021.
- H. Keurti, H.-R. Pan, M. Besserve, B. F. Grewe, and B. Schölkopf. Homomorphism autoencoder – learning group structured representations from observed transitions. In *International Conference on Machine Learning*, 2023.
- I. Khemakhem, D. Kingma, R. Monti, and A. Hyvarinen. Variational autoencoders and nonlinear ica: A unifying framework. In *International Conference on Artificial Intelligence and Statistics*, 2020a.
- I. Khemakhem, R. Monti, D. Kingma, and A. Hyvarinen. Ice-beem: Identifiable conditional energy-based deep models based on nonlinear ica. In *Advances in Neural Information Processing Systems*, 2020b.
- D. P. Kingma and J. Ba. Adam: A method for stochastic optimization. In *International Conference on Learning Representations*, 2015.

- D. P. Kingma and M. Welling. Auto-encoding variational bayes. In *International Conference on Learning Representations*, 2014.
- B. Kivva, G. Rajendran, P. K. Ravikumar, and B. Aragam. Identifiability of deep generative models without auxiliary information. In *Advances in Neural Information Processing Systems*, 2022.
- D. A. Klindt, L. Schott, Y. Sharma, I. Ustuzhaninov, W. Brendel, M. Bethge, and D. M. Paiton. Towards nonlinear disentanglement in natural data with temporal sparse coding. In *International Conference on Learning Representations*, 2021.
- M. Kocaoglu, C. Snyder, A. G. Dimakis, and S. Vishwanath. CausalGAN: Learning causal implicit generative models with adversarial training. In *International Conference on Learning Representations*, 2018.
- S. Lachapelle, Rodriguez Lopez, P., Y. Sharma, K. E. Everett, R. Le Priol, A. Lacoste, and S. Lacoste-Julien. Disentanglement via mechanism sparsity regularization: A new principle for nonlinear ICA. In *Conference on Causal Learning and Reasoning*, 2022.
- S. Lachapelle, T. Deleu, D. Mahajan, I. Mitliagkas, Y. Bengio, S. Lacoste-Julien, and Q. Bertrand. Synergies between disentanglement and sparsity: Generalization and identifiability in multi-task learning. In *International Conference on Machine Learning*, 2023a.
- S. Lachapelle, D. Mahajan, I. Mitliagkas, and S. Lacoste-Julien. Additive decoders for latent variables identification and cartesian-product extrapolation. In *Advances in Neural Information Processing Systems*, 2023b.
- T. Leemann, M. Kirchhof, Y. Rong, E. Kasneci, and G. Kasneci. When are post-hoc conceptual explanations identifiable? *Conference on Uncertainty in Artificial Intelligence*, 2023.
- A. Lei, B. Schölkopf, and I. Posner. Variational causal dynamics: Discovering modular world models from interventions. *Transactions on Machine Learning Research*, 2023.
- W. Liang, A. Kekić, J. von Kügelgen, S. Buchholz, M. Besserve, L. Gresele, and B. Schölkopf. Causal component analysis. In *Advances in Neural Information Processing Systems*, 2023.
- P. Lippe, S. Magliacane, S. Löwe, Y. M. Asano, T. Cohen, and E. Gavves. CITRIS: Causal identifiability from temporal intervened sequences. *arXiv preprint arXiv:2202.03169*, 2022.
- P. Lippe, S. Magliacane, S. Löwe, Y. M. Asano, T. Cohen, and E. Gavves. BISCUIT: Causal representation learning from binary interactions. In *Conference on Uncertainty in Artificial Intelligence*, 2023a.
- P. Lippe, S. Magliacane, S. Löwe, Y. M. Asano, T. Cohen, and E. Gavves. iCITRIS: Causal representation learning for instantaneous temporal effects. In *International Conference on Learning Representations*, 2023b.
- Y. Liu, Z. Zhang, D. Gong, M. Gong, B. Huang, A. van den Hengel, K. Zhang, and J. Qinfeng Shi. Identifying weight-variant latent causal models. *arXiv preprint arXiv:2208.14153*, 2023.

- F. Locatello, S. Bauer, M. Lucic, G. Raetsch, S. Gelly, B. Schölkopf, and O. Bachem. Challenging common assumptions in the unsupervised learning of disentangled representations. In *International Conference on Machine Learning*, 2019.
- F. Locatello, B. Poole, G. Raetsch, B. Schölkopf, O. Bachem, and M. Tschannen. Weakly-supervised disentanglement without compromises. In *International Conference on Machine Learning*, 2020.
- R. Lopez, N. Tagasovska, S. Ra, K. Cho, J. K. Pritchard, and A. Regev. Learning causal representations of single cells via sparse mechanism shift modeling. *Conference on Causal Learning and Reasoning*, 2023.
- K. Madan, N. R. Ke, A. Goyal, B. Schölkopf, and Y. Bengio. Fast and slow learning of recurrent independent mechanisms. In *International Conference on Learning Representations*, 2021.
- C. J. Maddison, A. Mnih, and Y. W. Teh. The concrete distribution: A continuous relaxation of discrete random variables. *International Conference on Machine Learning*, 2017.
- A. Mansouri, J. Hartford, K. Ahuja, and Y. Bengio. Object-centric causal representation learning. In *NeurIPS Workshop on Symmetry and Geometry in Neural Representations*, 2022.
- J. M. Mooij, S. Magliacane, and T. Claassen. Joint causal inference from multiple contexts. *Journal of Machine Learning Research*, 2020.
- G. Elyse Moran, D. Sridhar, Y. Wang, and D. Blei. Identifiable deep generative models via sparse decoding. *Transactions on Machine Learning Research*, 2022.
- H. Morioka and A. Hyvarinen. Connectivity-contrastive learning: Combining causal discovery and representation learning for multimodal data. In *International Conference on Artificial Intelligence and Statistics*, 2023.
- H. Morioka, H. Hälvä, and A. Hyvärinen. Independent innovation analysis for nonlinear vector autoregressive process. In *International Conference on Artificial Intelligence and Statistics*, 2021.
- J. R. Munkres. *Topology*. Prentice Hall, Inc., 2 edition, 2000.
- J.R. Munkres. *Analysis On Manifolds*. Basic Books, 1991.
- S. Nair, Y. Zhu, S. Savarese, and L. Fei-Fei. Causal induction from visual observations for goal directed tasks. *arXiv preprint arXiv:1910.01751*, 2019.
- R. Pamfil, N. Sriwattanaworachai, S. Desai, P. Pilgerstorfer, P. Beaumont, K. Georgatzis, and B. Aragam. Dynotears: Structure learning from time-series data. *International Conference on Artificial Intelligence and Statistics*, 2020.
- J. Pearl. *Causality*. Cambridge university press, 2009.
- J. Pearl. The seven tools of causal inference, with reflections on machine learning. *Commun. ACM*, 2019.

- J. Peters, J. M. Mooij, D. Janzing, and B. Schölkopf. Causal discovery with continuous additive noise models. *Journal of Machine Learning Research*, 2014.
- J. Peters, D. Janzing, and B. Schölkopf. *Elements of Causal Inference - Foundations and Learning Algorithms*. MIT Press, 2017.
- D. Pollard. *A User’s Guide to Measure Theoretic Probability*. Cambridge University Press, 2001.
- Jakob Runge, Peer Nowack, Marlene Kretschmer, Seth Flaxman, and Dino Sejdinovic. Detecting and quantifying causal associations in large nonlinear time series datasets. *Science Advances*, 2019.
- A. Schell and H. Oberhauser. Nonlinear independent component analysis for discrete-time and continuous-time signals. *The Annals of Statistics*, 2023.
- B. Schölkopf. Causality for machine learning. *arXiv preprint arXiv:1911.10500*, 2019.
- B. Schölkopf, F. Locatello, S. Bauer, N. R. Ke, N. Kalchbrenner, A. Goyal, and Y. Bengio. Toward causal representation learning. *IEEE - Advances in Machine Learning and Deep Neural Networks*, 2021.
- X. Shen, F. Liu, H. Dong, Q. Lian, Z. Chen, and T. Zhang. Weakly supervised disentangled generative causal representation learning. *Journal of Machine Learning Research*, 2022.
- S. Shimizu, P.O. Hoyer, A. Hyvärinen, and A. Kerminen. A linear non-gaussian acyclic model for causal discovery. *Journal of Machine Learning Research*, 2006.
- P. Sorrenson, C. Rother, and U. Köthe. Disentanglement by nonlinear ica with general incompressible-flow networks (GIN). In *International Conference on Learning Representations*, 2020.
- C. Squires, Y. Wang, and C. Uhler. Permutation-based causal structure learning with unknown intervention targets. *Conference on Uncertainty in Artificial Intelligence*, 2020.
- C. Squires, A. Seigal, S. Bhate, and C. Uhler. Linear causal disentanglement via interventions. In *International Conference on Machine Learning*, 2023.
- A. Taleb and C. Jutten. Source separation in post-nonlinear mixtures. *IEEE Transactions on Signal Processing*, 1999.
- V. Thomas, J. Pondard, E. Bengio, M. Sarfati, P. Beaudoin, M.-J. Meurs, J. Pineau, D. Precup, and Y. Bengio. Independently controllable factors. *arXiv preprint arXiv:1708.01289*, 2017.
- L. Tong, V.C. Soon, Y.F. Huang, and R. Liu. AMUSE: a new blind identification algorithm. In *IEEE International Symposium on Circuits and Systems*, 1990.
- L. Tong, Y. Inouye, and R.-w. Liu. Waveform-preserving blind estimation of multiple independent sources. *IEEE Transactions on Signal Processing*, 1993.
- B. Varici, E. Acarturk, K. Shanmugam, A. Kumar, and A. Tajer. Score-based causal representation learning from interventions: Nonparametric identifiability. In *Causal Representation Learning Workshop at NeurIPS 2023*, 2023a.

- B. Varici, E. Acartürk, K. Shanmugam, and A. Tajer. General identifiability and achievability for causal representation learning. *arXiv preprint arXiv:2310.15450*, 2023b.
- S. Volodin. CauseOccam : Learning interpretable abstract representations in reinforcement learning environments via model sparsity. Master’s thesis, École Polytechnique Fédérale de Lausanne, 2021.
- J. Von Kügelgen, Y. Sharma, L. Gresele, W. Brendel, B. Schölkopf, M. Besserve, and F. Locatello. Self-supervised learning with data augmentations provably isolates content from style. In *Advances in Neural Information Processing Systems*, 2021.
- J. von Kügelgen, M. Besserve, W. Liang, L. Gresele, A. Kekić, E. Bareinboim, D. Blei, and B. Schölkopf. Nonparametric identifiability of causal representations from unknown interventions. In *Advances in Neural Information Processing Systems*, 2023.
- M. J. Wainwright and M. I. Jordan. Graphical models, exponential families, and variational inference. *Found. Trends Mach. Learn.*, 2008.
- Q. Xi and B. Bloem-Reddy. Indeterminacy in generative models: Characterization and strong identifiability. In *International Conference on Artificial Intelligence and Statistics*, 2023.
- M. Yang, F. Liu, Z. Chen, X. Shen, J. Hao, and J. Wang. CausalVAE: Disentangled representation learning via neural structural causal models. In *IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 2021.
- D. Yao, D. Xu, S. Lachapelle, S. Magliacane, P. Taslakian, G. Martius, J. von Kügelgen, and F. Locatello. Multi-view causal representation learning with partial observability. *arXiv preprint arXiv:2311.04056*, 2023.
- W. Yao, G. Chen, and K. Zhang. Temporally disentangled representation learning. In *Advances in Neural Information Processing Systems*, 2022a.
- W. Yao, Y. Sun, A. Ho, C. Sun, and K. Zhang. Learning temporally causal latent processes from general temporal data. In *International Conference on Learning Representations*, 2022b.
- J. Zhang, K. Greenewald, C. Squires, A. Srivastava, K. Shanmugam, and C. Uhler. Identifiability guarantees for causal disentanglement from soft interventions. In *Advances in Neural Information Processing Systems*, 2023.
- X. Zheng, B. Aragam, P.K. Ravikumar, and E.P. Xing. DAGs with NO TEARS: Continuous optimization for structure learning. In *Advances in Neural Information Processing Systems*, 2018.
- Y. Zheng, I. Ng, and K. Zhang. On the identifiability of nonlinear ICA: Sparsity and beyond. In *Advances in Neural Information Processing Systems*, 2022.