

Intervening to Learn and Compose Causally Disentangled Representations

Alex Markham

University of Copenhagen

ALEX.MARKHAM@CAUSAL.DEV

Isaac Hirsch[†]

Jeri A. Chang[†]

Liam Solus

KTH Royal Institute of Technology

Bryon Aragam[†]

[†]University of Chicago

Editors: Bijan Mazaheri and Niels Richard Hansen

Abstract

In designing generative models, it is commonly believed that in order to learn useful latent structure, we face a fundamental tension between expressivity and structure. In this paper we challenge this view by proposing a new approach to training arbitrarily expressive generative models that simultaneously learn causally disentangled concepts. This is accomplished by adding a simple *context module* to an arbitrarily complex black-box model, which learns to process concept information by implicitly inverting linear representations from the model’s encoder. Inspired by the notion of intervention in a causal model, our module selectively modifies its architecture during training, allowing it to learn a compact joint model over different contexts. We show how adding this module leads to causally disentangled representations that can be composed for out-of-distribution generation on both real and simulated data. The resulting models can be trained end-to-end or fine-tuned from pre-trained models. To further validate our proposed approach, we prove a new identifiability result that extends existing work on identifying structured representations.

Keywords: causal disentanglement; generative models; compositional generalization; concept learning; interventions.

1. Introduction

Generative models have transformed information processing and demonstrated remarkable capacities for creativity in a variety of tasks ranging from vision to language to audio. The success of these models has been largely driven by modular, differentiable architectures based on deep neural networks that learn useful representations for downstream tasks. Recent years have seen increased interest in understanding and exploring the representations produced by these models through evolving lines of work on structured representation learning, identifiability and interpretability, disentanglement, and causal generative models. This work is motivated in part by the desire to produce performant generative models that also capture meaningful, semantic latent spaces that enable out-of-distribution (OOD) generation.

A key driver behind this work is the tradeoff between flexibility and structure, or expressivity and interpretability: Conventional wisdom suggests that to learn structured, interpretable representations, model capacity must be constrained, sacrificing flexibility and expressivity. This intuition is supported

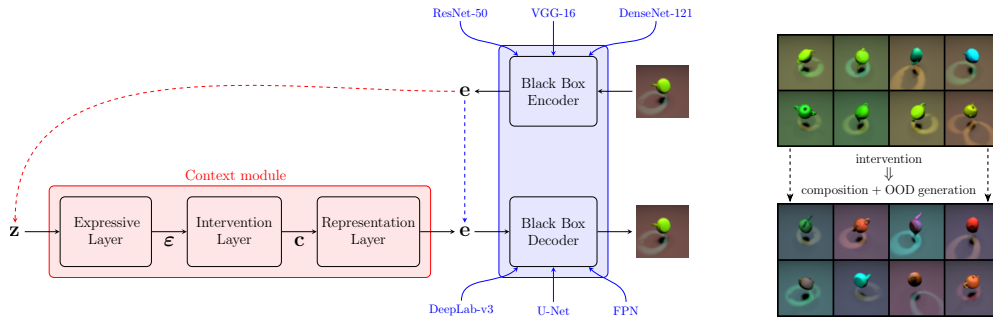


Figure 1: Overview (see (4) for notation map). Given a black-box encoder-decoder architecture (blue), we propose to prepend a *context module* (red) to the decoder. (left) Instead of passing the output of the encoder directly to the decoder (blue box + blue arrow), the embeddings e are passed through the context module, consisting of three distinct layers. The output of this module is then passed into the decoder. (right) The model learns to compose different concepts OOD, e.g., object and background colour, to values that never appear together in the training data.

by a growing body of work on nonlinear ICA (Hyvärinen and Pajunen, 1999), disentanglement (Bengio, 2013), and causal representation learning (Schölkopf et al., 2021). On the practical side, methods that have been developed to learn structured latent spaces tend to be bespoke to specific data types and models, and typically impose limitations on the flexibility and expressivity of the underlying models. Moreover, the most successful methods for learning structured representations typically impose fixed, known structure *a priori*, as opposed to learning this structure from data.

At the same time, there is also a growing body of work that suggests generative models already learn surprisingly structured latent spaces (e.g. Mikolov et al., 2013; Szegedy et al., 2013; Radford et al., 2015). This in turn suggests that existing models are already “close” to capturing useful structure, and perhaps only small modifications are needed. Since we already have performant models that achieve state-of-the-art results in generation and prediction, do we need to re-invent the wheel to achieve these goals? Our hypothesis is that we should be able to leverage the expressiveness of these models to build new models that learn latent structure from scratch. We emphasize that our goal is not to explain or interpolate the latent space of a pre-trained model, but rather to leverage known architectures to train a *new* model, either end-to-end or by fine-tuning an existing model.

In this paper, we adopt this perspective: We start with an arbitrarily complex black-box model, upon which we make no assumptions, then augment this model to learn causally disentangled representations. The idea is that the black-box model is already known to perform well on downstream tasks (e.g., generation, prediction), and so is flexible enough to capture complex patterns in data. We introduce a modular, end-to-end differentiable architecture for learning causally disentangled representations (Bengio, 2013; Schölkopf et al., 2021) that augments the original model with a simple module that retains its existing capabilities while enabling concept identification and intervention as well as composing together multiple concept interventions for OOD generation (Figure 1). The resulting latent structure is learned entirely from data, with no fixed structure imposed *a priori*. The motivation is to provide a framework for taking existing performant architectures and to train a new model that performs just as well, but with the added benefit of learning structured representations without imposing specific prior knowledge or structure. Since the model, unlike previous approaches, does not impose rigid latent structure, we theoretically validate the approach

with a novel identifiability result for structured representations recovered by the model. For a detailed discussion of related work, see Appendix A.

Contributions More specifically, we make the following contributions:

1. We introduce a *context module* that is attached to a decoder to learn concept representations by splitting each concept into a tensor slice, where each slice represents an interventional context in a reduced form structural equation model (SEM). This module can be attached to any decoder and trained end-to-end or fine-tuned (Figure 1).
2. We evaluate this module through quantitative and qualitative experiments on OOD generation, which is distinct from and significantly more challenging than OOD reconstruction as in prior work (e.g. Xu et al., 2022; Mo et al., 2024; Montero et al., 2022; Schott et al., 2022). We perform experiments on OOD generation on several real datasets as well as carefully controlled simulations.
3. We illustrate the adaptivity of our module to different architectures (including NVAE, Vahdat and Kautz, 2020) and against multiple ablations that control model capacity, data contexts, training protocols, and hyperparameters, ensuring that any differences in performance are directly attributable to the context module.
4. We fine-tune existing, pre-trained models after attaching the context module, giving substantial computational savings compared to end-to-end training while improving the learned representations and lending OOD generation abilities to pre-trained models.

As a matter of independent interest, we prove an identifiability result under concept interventions (Theorem 4), and discuss how the proposed architecture can be interpreted as an approximation to this model. We also introduce *quad*, a simple simulated visual environment for evaluating causal disentanglement, compositional abilities, and OOD generation. Due to space limitations, full technical proofs, along with detailed literature review and experimental details, can be found in the appendix.

2. Preliminaries

Let $\mathbf{x}$ denote the observations, e.g., pixels in an image, and $\mathbf{z}$ denote hidden variables, e.g., latent variables that are to be inferred from the pixels. A typical generative model consists of a decoder $p_\theta(\mathbf{x} | \mathbf{z})$ and an encoder $q_\phi(\mathbf{z} | \mathbf{x})$. After specifying a prior $p_\theta(\mathbf{z})$, this defines a likelihood by

$$p_\theta(\mathbf{x}) = \int p_\theta(\mathbf{x} | \mathbf{z}) p_\theta(\mathbf{z}) d\mathbf{z}. \quad (1)$$

Both the decoder and encoder are specified by deep neural networks that are trained end-to-end via standard techniques (Kingma and Welling, 2014; Rezende et al., 2014; Rezende and Mohamed, 2015). In this work, we focus on variational autoencoders (VAEs) since they are often used to represent structured and semantic latent spaces, which is our focus. Unlike traditional structured VAEs, we do not impose a specific structure *a priori*.

Causal disentanglement Our goal is to learn distinct latent factors that can be composed together. This is closely related to causal disentanglement, which we recall informally here:

Definition 1 (Bengio, 2013; Thomas et al., 2017) *A causally disentangled latent space consists of separately controllable (intervenable) latent factors.*

In particular, we want to manipulate and intervene on latent factors in a causally meaningful way; crucially this allows us to compose distinct concepts together without needing to learn the full causal graph in the latent space. This notion of disentanglement is also different from the axis-alignment notion (Kim and Mnih, 2018; Chen et al., 2018). The “causal” aspect comes from interpreting “separately controllable” to mean “able to intervene on independent causal mechanisms”, as suggested by Locatello et al. (2019); see also Schölkopf et al. (2021). See Appendix B for formal details. Rather than attempt to reconstruct the full causal DAG model (as in Yang et al., 2021; Squires et al., 2023; Buchholz et al., 2023), we aim for the more practical goal of learning different interventional contexts through distributional invariances; thus our approach strikes a balance between easy-to-learn unstructured latent spaces and hard-to-learn but desirable fully structured causal DAG latent spaces.

Measuring causal disentanglement A standard approach to quantifying causal disentanglement is to evaluate the recovery of the latent causal graph; since our explicit goal is to avoid learning this graph, we propose a different metric that is more suitable to downstream applications. To measure causal disentanglement, we use OOD generation on novel interventions as a metric. We train models so that they never see examples of certain latent factors composed together (equivalently, multiple simultaneous interventions on distinct factors), so these examples are genuinely OOD. By withholding these examples during training and building a model that is capable of intervening on multiple factors, we can compare the quality of OOD samples generated from our model to the held-out OOD samples. In our experiments, we use the sliced Wasserstein (SW) metric (18); see Appendix D.1 for details.

OOD reconstruction vs. OOD generation Previous work has evaluated the OOD capabilities of generative models by using OOD reconstruction as a metric (e.g. Xu et al., 2022; Mo et al., 2024; Montero et al., 2022; Schott et al., 2022). There is a crucial difference between OOD reconstruction and OOD generation: By OOD generation, we mean the ability to conditionally generate and combine novel interventions on existing concepts. For example, suppose during training the model only sees small, red objects along with big, blue objects. At test time, we wish to generate OOD samples of big, red objects or small, blue objects that were never seen during training.

The crucial difference is that while we can *always* attempt to reconstruct *any* sample (i.e., OOD or not) and evaluate its reconstruction error, it is not always possible to have the model *generate* a new OOD sample. Unless the model learns specific latent factors corresponding to concepts such as size or colour, we cannot control (i.e., intervene) the size or colour of random samples from the model.

Thus, OOD generation not only evaluates the ability of a model to compose learned concepts in new ways, but also the ability of a model to identify and capture underlying concepts of interest. Accordingly, we argue that OOD generation is more appropriate to evaluate causal disentanglement since it captures *both* the model’s ability to learn concepts as distinct latent factors *as well as* compose them together in novel ways.

Set-up and approach Our set-up is the following: We have a black-box encoder and decoder, on which we make no assumptions other than the encoder outputs a latent code corresponding to $\mathbf{x}$. In the notation of (1), this corresponds to $\mathbf{z}$, however, to distinguish the black-box embeddings from our model, we will denote the black-box embeddings hereafter by $\mathbf{e}$. Our plan is to work solely with the embeddings $\mathbf{e}$ and learn how to extract linear concept representations from $\mathbf{e}$ such that distinct concepts can be intervened upon and composed together to create novel, OOD samples.

Our approach is motivated by the following incongruence: On the one hand, it is well-established that the latent spaces of generative models are typically entangled and semantically misaligned, fail to generalize OOD, and suffer from posterior collapse (which has been tied to latent variable nonidentifiability, Wang et al., 2021). On the other hand, they still learn structured latent spaces that can be interpolated, are “nearly” identifiable (Willettts and Paige, 2021; Reizinger et al., 2022), and represent abstract concepts linearly (Mikolov et al., 2013; Szegedy et al., 2013; Radford et al., 2015). So, to learn latent structure, we simply need to extract the right concept representations from the already performant embeddings of existing architectures and learn how to compose them together.

3. Architecture

We start with a black-box encoder-decoder architecture along with d_c concepts of interest, denoted by $\mathbf{c} = (\mathbf{c}_1, \dots, \mathbf{c}_{d_c})$. Our objectives are three-fold:

1. To extract concepts from black-box embeddings $\mathbf{e}$ as linear projections $C\mathbf{e}$;
2. To compose concepts together in a single, transparent model that captures how different concepts are related;
3. To accomplish both of these objectives without restricting the black-box encoder or decoder architectures, and without sacrificing sampling quality.

A key intuition is that composition can be interpreted as a type of intervention in the latent space. This is sensible since interventions in a causal model are a type of nontrivial distribution shift. Thus, the problem takes on a causal flavour which we leverage to build our architecture. The difficulty with this from the causal modeling perspective is that encoding structural assignments and/or causal mechanisms directly into a feed-forward neural network is tricky, because edges between nodes within the same layer are not allowed in a feed-forward network but are required for the usual DAG representation of a causal model. To bypass this, we use the *reduced form* of a causal model which has a clear representation as a bipartite directed graph, with arrows only from exogenous to endogenous variables, making it conducive to being embedded into a neural network. The tradeoff is that this does not learn a causal DAG, however it still enables the model to perform interventions directly in the latent space and to leverage invariances between interventional contexts.

3.1. Overview

Before outlining the architectural details, we provide a high-level overview of the main idea. A traditional decoder transforms the embeddings $\mathbf{e}$ into the observed variables $\mathbf{x}$, and the encoder operates in reverse by encoding $\mathbf{x}$ into $\mathbf{e}$. Thus,

$$\mathbf{x} \xrightarrow{\text{encoder}} \mathbf{e} \xrightarrow{\text{decoder}} \hat{\mathbf{x}}.$$

Based on an extensive body of empirical work that shows concepts are linearly represented (Mikolov et al., 2013; Szegedy et al., 2013; Radford et al., 2015, see Appendix A for more discussion), we represent concepts as linear projections of the embeddings $\mathbf{e}$: Each concept $\mathbf{c}_j$ can be approximated as $\mathbf{c}_j \approx C_j \mathbf{e}$. This is modeled via a linear layer $\mathbf{c} \rightarrow \mathbf{e}$ that implicitly inverts this relationship in the decoder.

We further model the relationships between concepts with a linear SEM:

$$\mathbf{c}_j = \sum_{k=1}^{d_c} \alpha_{kj} \mathbf{c}_k + \varepsilon_j, \quad \alpha_{kj} \in \mathbb{R}. \quad (2)$$

By reducing this SEM and solving for $\mathbf{c}$, we deduce that

$$\mathbf{c} = A_0 \varepsilon, \text{ where } \begin{cases} \mathbf{c} = (\mathbf{c}_1, \dots, \mathbf{c}_{d_c}) \\ \varepsilon = (\varepsilon_1, \dots, \varepsilon_{d_c}) \end{cases}. \quad (3)$$

This is known as the *reduced form* of the SEM (2). We model this with a linear layer $\varepsilon \rightarrow \mathbf{c}$, where the weights in this layer correspond to the matrix A_0 . This layer will be used to encode the SEM between the concepts, which will be used to implement causal interventions.

Remark 2 *Due to the reduced form SEM above, our approach does not and cannot model the structural causal model encoded by the α_{kj} . What is important is that the reduced form $\mathbf{c} = A_0 \varepsilon$ still encodes causal invariances and interventions, which is enough in our setting, without directly estimating a causal graph.*

In principle, since the exogenous variables ε_j are independent, we could treat ε as the input latent space to the generative model. Doing this, however, incurs two costs: 1) To conform to standard practice, ε would have to follow an isotropic Gaussian prior, and 2) It enforces artificial constraints on the latent dimension $\dim(\varepsilon)$. For this reason, we use a second expressive layer $\mathbf{z} \rightarrow \varepsilon$ that gradually transforms $\mathbf{z} \sim \mathcal{N}(0, I)$ into ε . This allows for $\dim(\mathbf{z})$ to be larger and more expressive than d_c (in practice, we set $\dim(\mathbf{z})$ to be a multiple of $\dim(\varepsilon)$), and for ε to be potentially non-Gaussian.

The final decoder architecture can be decomposed at a high-level as follows (Figure 1):

$$\mathbf{z} \longrightarrow \varepsilon \longrightarrow \mathbf{c} \longrightarrow \mathbf{e} \longrightarrow \mathbf{x}. \quad (4)$$

Implementation details for each of these layers can be found in the next section.

3.2. Details

We assume given a black-box encoder-decoder pair, denoted by $\text{enc}(\mathbf{x})$ and $\text{dec}(\mathbf{z})$, respectively. We modify the decoder by appending a *context module* to the decoder. The context module has three layers: A *representation layer*, an *intervention layer*, and an *expressive layer*.

1. The first *representation layer*, as the name suggests, learns to represent concepts by implicitly inverting their linear representations $\mathbf{c}_j = C_j \mathbf{e}$ from the embeddings of the encoder. This is a linear layer between $\mathbf{c} \rightarrow \mathbf{e}$ in (4).
2. The *intervention layer* embeds these concept representations into the reduced form SEM (3). Each concept has its own context where it has been intervened upon. A key part of this layer is how it can be used to learn and enforce interventional semantics through this shared SEM. This layer corresponds to the second layer $\varepsilon \rightarrow \mathbf{c}$ in (4).
3. The *expressive layer* is used to reduce the (potentially very large) input latent dimension of independent Gaussian inputs down to a smaller space of non-Gaussian exogenous noise variables for the intervention layer, corresponding to the first layer $\mathbf{z} \rightarrow \varepsilon$ in (4). This is implemented as independent, deep MLPs that gradually reduce the dimension in each layer: If it is desirable to preserve Gaussianity for the exogenous variables, linear activations can be used in place of nonlinear activations (e.g., ReLU).

Because the intervention layer is modeled after an SEM, it is straightforward to perform latent concept interventions using the calculus of interventions in an SEM.

Remark 3 *Crucially, no part of this SEM is fixed or known—everything is trained end-to-end. In particular, we do not assume a known causal DAG or even a known causal order. This stands in contrast to previous work on causal generative models (e.g. Kocaoglu et al., 2017; Yang et al., 2021).*

Concept interventions Assume for now that the concepts are each one-dimensional with $\dim(\mathbf{c}_j) = \dim(\boldsymbol{\varepsilon}_j) = 1$; extensions to multi-dimensional concepts are straightforward and explained below. The structural coefficients $\alpha_{kj} \in \mathbb{R}$ capture direct causal effects between concepts, with $\alpha_{kj} \neq 0$ indicating the presence of an edge $\mathbf{c}_k \rightarrow \mathbf{c}_j$. An intervention on the j th concept requires deleting each incoming edge; i.e., setting $\alpha_{\cdot j} = 0$, updating the outgoing edges from $\boldsymbol{\varepsilon}_j$, as well as replacing $\boldsymbol{\varepsilon}_j$ with a new $\boldsymbol{\varepsilon}'_j$, which results from updating the incoming edges to $\boldsymbol{\varepsilon}_j$ in the preceding layer. Thus, during training, when we intervene on $\mathbf{c}_j$, we zero out the row $\alpha_{\cdot j}$ while replacing the column $\alpha_{j\cdot}$ with a new column β_j that is trained specifically for this interventional setting. The result is $d_c + 1$ intervention-specific layers: A_0 for the observational setting, and $A_1, \dots, A_{d_c}$ for each interventional setting. This yields a three-way tensor where each slice of this tensor captures a different context that corresponds to intervening on different concepts. At inference time, to generate interventional samples we simply swap out the A_0 for A_j . Moreover, multi-target interventions can be sampled by zeroing out multiple rows and substituting multiple β_j ’s and $\boldsymbol{\varepsilon}_j$ ’s.

Multi-dimensional concepts Everything above goes through if we allow additional flexibility with $\dim(\mathbf{c}_j) > 1$ and $\dim(\boldsymbol{\varepsilon}_j) > 1$, potentially even with $\dim(\mathbf{c}_j) \neq \dim(\boldsymbol{\varepsilon}_j)$. In practice, we implement this via two width parameters $w_\varepsilon = \dim(\boldsymbol{\varepsilon}_j)$ and $w_c = \dim(\mathbf{c}_j)$. The design choice of enforcing uniform dimensionality is not necessary, but is made here since in our experiments there did not seem to be substantial advantages to choosing nonuniform widths. With these modifications, α_{kj} becomes a $w_\varepsilon \times w_c$ matrix; instead of substituting rows and columns above, we now substitute slices in the obvious way. The three-way tensor A is now $(d_c + 1) \times w_\varepsilon d_c \times w_c d_c$ -dimensional.

Identifiability An appealing aspect of this architecture is that it can be viewed as an approximation to an identifiable model over causally disentangled concepts. Formally, we assume given observations $\mathbf{x}$, concepts $\mathbf{c}$, and a collection of embeddings $\mathbf{e}$ such that $\mathbf{x} = f(\mathbf{e})$. Additionally assume a noisy version of the linear representation hypothesis, i.e., $\mathbf{c}_j = C_j \mathbf{e} + \boldsymbol{\varepsilon}_j$ with $\boldsymbol{\varepsilon}_j \sim \mathcal{N}(0, \Omega_j)$. Then we have the following identifiability result:

Theorem 4 (Identifiability) *Assume that the rows of each C_j come from a linearly independent set and f is injective and differentiable. Then, given single-node interventions on each concept $\mathbf{c}_j$, the representations C_j and latent concept distribution $p(\mathbf{c})$ are identifiable.*

For a formal statement, see Theorem 5 in Appendix B. As opposed to solving a causal representation learning problem, the causal semantics used in this paper have been ported into a generative model that give solutions to the problem of learning latent *distribution* structure, as in Theorem 4. We note also that the proof of Theorem 4 itself is nontrivial, and involves tracing the effects of interventions through the linear representations; since our focus is on implementation and practical aspects, we defer further discussion to Appendix B.

3.3. Summary of architecture

The context module thus proposed offers the following appealing desiderata in practice:

1. **Black-box.** There is no coupling between the embedding dimension and the number of concepts, or the complexity of the SEM. This allows for *arbitrary* black-box encoder-decoder architectures to be used to learn embeddings, from which a causal model is then trained on top of. This can be trained end-to-end, fine-tuned after pre-training, or even partially frozen, allowing for substantial computational savings.
2. **Causality.** The intervention layer is a genuine causal model that provides rigorous causal semantics that allow sampling from arbitrary concept interventions, including interventions that have not been seen during training.
3. **Identifiability.** The architecture is based on an approximation to an identifiable concept model (Theorem 4), which provides formal justification for the intervention layer as well as reproducibility assurances.
4. **Flexibility.** The context module is itself arbitrarily flexible, so there is no risk of information loss in representing concepts with the embeddings $\mathbf{e}$. Of course, information loss is possible if we compress this layer too much (e.g., by choosing d_c , k , w_ε , or w_c too small), but this is a design choice and not a constraint of the module itself.

Thus the only tradeoff between representational capacity and causal semantics is design-based: The architecture itself imposes no constraints. The causal model can be arbitrarily flexible and chosen independent of the encoder-decoder pair, which are also allowed to be arbitrary.

4. Empirical evaluation

Open-source Python code and instructions for reproducing the results are available at <https://github.com/Alex-Markham/context-module>. We evaluate the proposed context module on four different tasks:

1. **Concept learning** (Section 4.1): *How well does the context module learn representations that correspond to each concept?* To measure this, we generate interventional samples by intervening on a single concept, and directly compare these interventional samples to held-out samples using the sliced Wasserstein (SW) metric (18).
2. **Composition** (Section 4.2): *How well does the context module learn to compose concepts together in novel ways that have not been seen during training?* To measure this, we generate samples by intervening on multiple concepts, and directly compare these new samples to held-out samples using the SW metric. By design, interventions on multiple concepts are never seen during training, so this is genuinely OOD.
3. **Reconstruction** (Section 4.3): *How well does the context module reconstruct samples from each context?* This can be evaluated using standard measures of generation quality; we use bits per dimension (BPD) for comparison with prior work. Since this is computed across all contexts, this captures the model’s ability to reconstruct samples from diverse contexts.
4. **Fine-tuning** (Section 4.4): *Can pre-trained models leverage the context module via fine-tuning to improve their learned representations and enable compositional generation?* Instead of training end-to-end, we start with a pre-trained model and fine-tune it with the context module attached. This is evaluated on the concept learning, composition, and reconstruction tasks.

For concept learning and composition, including OOD generation (Section 2), we report the sliced Wasserstein (SW) distance (Bonneel et al., 2015; Flamary et al., 2021) between a held-out sample from the ground-truth distribution and a sample generated (not reconstructed) by the trained model. For reconstruction, we report the standard ELBO loss (Kingma and Welling, 2014) in bits per dimension.

Methods, ablations, baselines We implement three variants of the proposed context module: 1) Full end-to-end training; 2) Fine-tuning all weights from a pre-trained base model; 3) Fine-tuning only the context module weights while leaving the pre-trained base model frozen. We evaluate our approach by augmenting two base architectures—a lightweight convolutional β -VAE (Higgins et al., 2017) and a deep hierarchical network, NVAE (Vahdat and Kautz, 2020)—with our context module (CM). This allows us to compare directly to ablations using the base architectures without the context module: 1) **base**, which uses only observational data; 2) **base pooled**, which pools all contexts together into a single dataset; and 3) **CM pooled**, which pools all contexts together into a single context using the context module. Thus, for each dataset, every aspect of the ablation is controlled: architecture, hyperparameters, random seeds, training protocol, etc. The only difference is whether or not the context module is inserted and how many contexts are used. We can therefore distinguish performance differences due to data diversity and/or architectural complexity. We additionally include comparisons to three other methods: Ada-GVAE (Locatello et al., 2020), β TCVAE (Chen et al., 2018), and TVAE (Ishfaq et al., 2018).

Datasets We compare these models to our context module on four datasets: 3DIdent (Zimmermann et al., 2021), Morpho MNIST (Castro et al., 2019), and two versions of a new semi-synthetic environment called *quad*—one with independent concepts and another with (causally) dependent concepts. While 3DIdent and Morpho MNIST are established baselines, *quad* is a novel visual environment with intervenable latents (colour, shape, size, orientation) that explicitly enables sampling from ground-truth OOD composed contexts. See Appendix C for more details on this dataset. We introduce *quad* in order to carefully evaluate model performance in a controlled environment over different random seeds (and hence assess statistical variability). Thus, our evaluations use *quad* to assess reproducibility and statistical significance, whereas 3DIdent and MNIST are more appropriate for testing our approach on downstream tasks such as image generation.

Due to space constraints, we report only a representative slice of the experiments here; full details and additional experiments can be found in Appendices C–E. This includes detailed performance breakdowns across contexts; the effects of regularization and architectural capacity; and additional figures.

4.1. Concept learning

Table 1, in the **Concept** columns, shows a quantitative evaluation of concept learning on the *quad*, MNIST, and 3DIdent datasets. We draw three main conclusions from this table, by comparing our context module variants to each of the other methods (rows) in the ablations and baselines.

First, comparing the three **CM** rows to the ablations in Table 1 (specifically, **base** and **base pooled**), we see that the context module excels at learning individual concepts across all datasets and models, despite the base models achieving slightly better average reconstruction performance (see Section 4.3). Moreover, Figure 2 shows that there is essentially no perceptual difference between samples from the base models and the context module. Additionally, across multiple training runs

Table 1: Context module (CM) vs. fine-tuning, ablations, and baselines on three tasks across four datasets. Best performance per task/dataset in bold. Base model indicated by * = lightweight VAE, ** = NVAE.

		quad* (independent)			quad* (dependent)			MNIST**			3DIdent**		
		Concept	Compo	Recon	Concept	Compo	Recon	Concept	Compo	Recon	Concept	Compo	Recon
ours	CM (end-to-end)	0.063	0.081	0.519	0.116	0.127	0.828	0.035	-	0.135	0.014	0.051	0.528
	CM (fine-tune)	0.059	0.086	0.518	0.112	0.121	0.840	0.037	-	0.136	0.031	0.063	0.549
	CM (frozen)	0.086	0.132	0.559	0.121	0.134	0.889	0.052	-	0.161	0.048	0.063	0.609
ablations	base	0.175	0.288	0.444	0.170	0.222	0.767	0.102	-	0.271	0.056	0.092	4.966
	base pooled	0.166	0.248	0.446	0.167	0.212	0.746	0.105	-	0.139	0.049	0.056	0.559
	CM pooled	0.202	0.306	0.840	0.166	0.203	0.752	0.099	-	0.143	0.041	0.068	0.516
baselines	Ada-GVAE	0.335	0.355	2.017	0.375	0.377	2.443	0.500	-	1.035	0.249	0.251	3.267
	BetaTCVAE	0.330	0.346	1.973	0.366	0.374	2.408	0.484	-	1.036	0.257	0.254	3.261
	TVAE	0.334	0.341	2.048	0.370	0.368	2.455	0.486	-	1.037	0.254	0.256	3.282

over different random seeds, our context module consistently outperforms the base models across quad concepts—in this case the simpler model/data allows more replicates to be run to evaluate statistical significance and standard errors.

Second, comparing the **CM** rows to the **base** row of Table 1, we see that the context module outperforms the base model trained only on the observational (rather than pooled) data. This confirms that the context module is indeed successfully exploiting learned invariances across the different contexts, providing it with performance and computational resource advantages compared to models trained independently on the different contexts.

Third, comparing the **CM** rows to the ablation in **CM pooled**—which pools all the data into a single context in order to circumvent the context module’s intervention mechanism while leaving the rest of the context module architecture intact—we see that **CM** consistently outperforms this ablation. This demonstrates that the increased performance of the context module is due our module’s intervention mechanism and its resulting ability to better leverage context-separated data, as opposed to spuriously benefiting from a different architecture or increased complexity.

Finally, for a qualitative evaluation, we look at Figure 2, in the Concept 1 and Concept 2 rows, to see concept learning in CM+NVAE across three different datasets. For example, looking at the 3DIdent results in Figure 2 (middle), the Observation row shows the learned observational distribution of images to contain background colours and object colours both in the orange–green range (). The Concept 1 row, produced by intervening on the learned concept of background colour, shows that the background colour is shifted to the blue–red range while other features, like object colour, remain invariant with the observational samples. Analogously, the Concept 2 row shows a shift of object colour to the blue–red range () while the other features remain invariant. These illustrate that CM+NVAE successfully learned the concepts of background colour and object colour, in line with the notion of causal disentanglement in Definition 1.

4.2. Concept composition

Table 1, in the **Compo** columns, provides a quantitative evaluation of concept composition on the quad and 3DIdent datasets, while the Compo row of Figure 2 provides a qualitative evaluation across the MNIST and 3DIdent datasets. Since Morpho MNIST does not include ground truth compositional samples, we are not able to evaluate OOD generation on this dataset. In each evaluation, the model

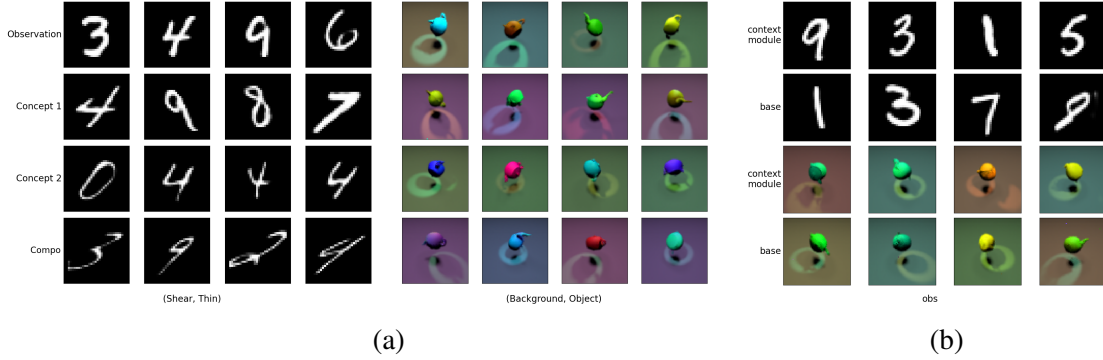


Figure 2: (a) Examples of concept learning and composition in MNIST (left) and 3DIdent (right) using CM (end-to-end), our context module augmenting the base NVAE—these images are *generated*, not reconstructed. The first row shows generated observational samples; The second and third rows show samples of *learned concepts*, generated by intervening on individual concepts, e.g., ‘Concept 1’ is ‘scaled’ for MNIST; The final row shows *OOD composition*, generated by simultaneously intervening on pairs of learned concepts. (b) Observational samples from the context module vs. the base model show no perceptual loss.

uses the intervention layer (Section 3.2) to intervene on pairs of learned concepts, i.e., to compose two concepts together. By holding out from training examples where concepts are composed, we are assured that the model has not seen any of these compositions. The only way to compose these concepts together is through the implicit causal model (via the reduced form SEM (3)) that is learned during training. We evaluate model performance with the SW metric (18), comparing concept compositions from the model against ground-truth hold out compositions.

The same patterns in performance hold here for concept composition as we saw for concept learning in Section 4.1—namely, our context module consistently achieves better performance than baselines and ablations, demonstrating successful concept composition due to the module’s intervention mechanism. For example, looking again at the 3DIdent results in Figure 2 (middle), we see that the colour range of both the background and the object are shifted from the orange–green range (orange to green) to the blue–red range (blue to red).

These experiments demonstrate the context module’s ability to learn causally disentangled representations through both concept learning and composition.

4.3. Reconstruction

The **Recon** columns of Table 1 compare the validation reconstruction performance of our approach against the alternatives. We draw two conclusions from these results. First, adding the context module incurs only a minor quantitative cost in reconstruction performance, while granting additional concept learning and composition abilities explored in Sections 4.1-4.2. Second, Figure 2 demonstrates that despite this small quantitative cost, the perceptual quality of the samples is still very high and indistinguishable from the **base** model. This validates the central motivation behind the context module, namely that it can imbue existing black-box models with additional concept learning and causal composition abilities *without* sacrificing perceptual metrics or downstream performance.

4.4. Fine-tuning

Finally, an important aspect of our method is that it need not be trained end-to-end but can instead be successfully fine-tuned from a frozen base mode. The second (**CM (fine-tune)**) and third (**CM (frozen)**) rows of Table 1 evaluate the context module on concept learning and composition using two fine-tuning variants. Both variants start with a pre-trained model and then insert the context module and resume training, either only updating the context module weights in the **frozen** variant, or also fine-tuning the non-context module weights in the **fine-tune** variant.

We see that the fine-tuned models consistently outperform both the **base** variants as well as the baselines, despite minor fluctuations in performance compared to end-to-end training (**CM (end-to-end)**). Hence, both fine-tuning variants offer computational savings (ranging from 2–10× in GPU hours) compared to the full end-to-end training while still substantially outperforming all other baselines and ablations in Table 1. Choosing between end-to-end training and the two variants allows for a tradeoff between computational cost and performance, all with improved representations and OOD generation abilities compared to the alternatives.

5. Conclusion and Limitations

Designing generative models with structured latent spaces that enable intervention, composition, and OOD generation is an outstanding challenge. Existing approaches either sacrifice structure for flexibility, or flexibility for structure. We argue that such a tradeoff is not necessary and that arbitrarily flexible models can indeed learn useful structure.

To this end, we propose a simple context module that can be appended to a black-box decoder that learns causally disentangled latent spaces. The context module enables genuinely OOD, compositional generation without hardcoding prior knowledge or structure into the model, and it avoids the difficult problem of latent causal discovery.

Our experiments provide both quantitative and qualitative evidence for this. While existing evaluations have focused on reconstruction as a metric to quantify OOD performance, we argue that OOD generation is in fact more appropriate since it also measures the ability of a model to learn useful concepts that can be manipulated at inference time.

Since our module imposes no hard constraints, it may be useful in combination with other modules. Understanding how the context module interacts with other architectures is a natural next step and a promising direction for future work. We also encourage the community to explore alternative metrics for evaluating genuine OOD generalization.

Finally, understanding the cost and tradeoffs of variational approximations compared to, say, exact maximum likelihood estimation is important, especially in causal applications. Indeed, this can make estimating the exact causal graph much more challenging. This motivates our fundamental design principle: to use a “causal enough” architecture to obtain composability and OOD generation while otherwise leveraging already empirically proven non-causal architectures.

ACKNOWLEDGMENTS

The work by AM was supported by Novo Nordisk Foundation Grant NNF20OC0062897. This research was supported in part through the computational resources and staff contributions provided for the Mercury high performance computing cluster at The University of Chicago Booth School of Business which is supported by the Office of the Dean.

REFERENCES

- Kartik Ahuja, Divyat Mahajan, Yixin Wang, and Yoshua Bengio. Interventional causal representation learning. In *Proceedings of the 40th International Conference on Machine Learning, ICML'23*. JMLR, 2023.
- Kartik Ahuja, Amin Mansouri, and Yixin Wang. Multi-domain causal representation learning via weak distributional invariances. In *International Conference on Artificial Intelligence and Statistics*, pages 865–873. PMLR, 2024.
- Carl Allen and Timothy Hospedales. Analogies explained: Towards understanding word embeddings. In *International Conference on Machine Learning*, pages 223–231. PMLR, 2019.
- Sanjeev Arora, Yuanzhi Li, Yingyu Liang, Tengyu Ma, and Andrej Risteski. A latent variable model approach to PMI-based word embeddings. *Transactions of the Association for Computational Linguistics*, 4:385–399, 2016.
- Matthew Ashman, Chao Ma, Agrin Hilmkil, Joel Jennings, and Cheng Zhang. Causal reasoning in the presence of latent confounders via neural ADMG learning. In *The Eleventh International Conference on Learning Representations*, 2022.
- David Bau, Bolei Zhou, Aditya Khosla, Aude Oliva, and Antonio Torralba. Network dissection: Quantifying interpretability of deep visual representations. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 6541–6549, 2017.
- Yoshua Bengio. Deep learning of representations: Looking forward. In *Statistical Language and Speech Processing: First International Conference, SLSP 2013, Tarragona, Spain, July 29-31, 2013. Proceedings 1*, pages 1–37. Springer, 2013.
- Simon Bing, Urmi Ninad, Jonas Wahl, and Jakob Runge. Identifying linearly-mixed causal representations from multi-node interventions. In *Causal Learning and Reasoning*, pages 843–867. PMLR, 2024.
- Nicolas Bonneel, Julien Rabin, Gabriel Peyré, and Hanspeter Pfister. Sliced and radon Wasserstein barycenters of measures. *Journal of Mathematical Imaging and Vision*, 51(1):22–45, 2015.
- Johann Brehmer, Pim De Haan, Phillip Lippe, and Taco S Cohen. Weakly supervised causal representation learning. *Advances in Neural Information Processing Systems*, 35:38319–38331, 2022.
- Simon Buchholz, Goutham Rajendran, Elan Rosenfeld, Bryon Aragam, Bernhard Schölkopf, and Pradeep Ravikumar. Learning linear causal representations from interventions under general nonlinear mixing. *arXiv preprint arXiv:2306.02235*, 2023.
- Collin Burns, Haotian Ye, Dan Klein, and Jacob Steinhardt. Discovering latent knowledge in language models without supervision. *arXiv preprint arXiv:2212.03827*, 2022.
- Daniel C Castro, Jeremy Tan, Bernhard Kainz, Ender Konukoglu, and Ben Glocker. Morpho-MNIST: Quantitative assessment and diagnostics for representation learning. *Journal of Machine Learning Research*, 20(178):1–29, 2019.

- Ricky TQ Chen, Xuechen Li, Roger B Grosse, and David K Duvenaud. Isolating sources of disentanglement in variational autoencoders. *Advances in neural information processing systems*, 31, 2018.
- Zhi Chen, Yijie Bei, and Cynthia Rudin. Concept whitening for interpretable image recognition. *Nature Machine Intelligence*, 2(12):772–782, 2020.
- Alexis Conneau, German Kruszewski, Guillaume Lample, Loïc Barrault, and Marco Baroni. What you can cram into a single vector: Probing sentence embeddings for linguistic properties. *arXiv preprint arXiv:1805.01070*, 2018.
- Hoagy Cunningham, Aidan Ewart, Logan Riggs, Robert Huben, and Lee Sharkey. Sparse autoencoders find highly interpretable features in language models. *ICLR 2024*, 2024.
- Muong Ding, Constantinos Daskalakis, and Soheil Feizi. GANs with conditional independence graphs: On subadditivity of probability divergences. In *International Conference on Artificial Intelligence and Statistics*, pages 3709–3717. PMLR, 2021.
- Nelson Elhage, Tristan Hume, Catherine Olsson, Nicholas Schiefer, Tom Henighan, Shauna Kravec, Zac Hatfield-Dodds, Robert Lasenby, Dawn Drain, Carol Chen, Roger Grosse, Sam McCandlish, Jared Kaplan, Dario Amodei, Martin Wattenberg, and Christopher Olah. Toy models of superposition. *Transformer Circuits Thread*, 2022.
- Jesse Engel, Matthew Hoffman, and Adam Roberts. Latent constraints: Learning to generate conditionally from unconditional generative models. *arXiv preprint arXiv:1711.05772*, 2017.
- Kawin Ethayarajh, David Duvenaud, and Graeme Hirst. Towards understanding linear word analogies. *arXiv preprint arXiv:1810.04882*, 2018.
- Rémi Flamary, Nicolas Courty, Alexandre Gramfort, Mokhtar Z. Alaya, Aurélie Boisbunon, Stanislas Chambon, Laetitia Chapel, Adrien Corenflos, Kilian Fatras, Nemo Fournier, Léo Gautheron, Nathalie T.H. Gayraud, Hicham Janati, Alain Rakotomamonjy, Ievgen Redko, Antoine Rolet, Antony Schutz, Vivien Seguy, Danica J. Sutherland, Romain Tavenard, Alexander Tong, and Titouan Vayer. Pot: Python optimal transport. *Journal of Machine Learning Research*, 22(78):1–8, 2021. URL <http://jmlr.org/papers/v22/20-451.html>.
- Hidde Fokkema, Tim van Erven, and Sara Magliacane. Sample-efficient learning of concepts with theoretical guarantees: From data to concepts without interventions. *arXiv preprint arXiv:2502.06536*, 02 2025. URL <https://arxiv.org/pdf/2502.06536.pdf>.
- Juan L. Gamella, Armeen Taeb, Christina Heinze-Deml, and Peter Bühlmann. Characterization and greedy learning of Gaussian structural causal models under unknown interventions. *arXiv preprint arXiv:2211.14897*, 2022.
- Mathieu Germain, Karol Gregor, Iain Murray, and Hugo Larochelle. MADE: Masked autoencoder for distribution estimation. In *International conference on machine learning*, pages 881–889. PMLR, 2015.

- Alex Gittens, Dimitris Achlioptas, and Michael W Mahoney. Skip-gram – zipf + uniform = vector additivity. In *Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 69–76, 2017.
- Luigi Gresele, Julius Von Kügelgen, Vincent Stimper, Bernhard Schölkopf, and Michel Besserve. Independent mechanism analysis, a new concept? *Advances in Neural Information Processing Systems*, 34, 2021.
- Wes Gurnee, Neel Nanda, Matthew Pauly, Katherine Harvey, Dmitrii Troitskii, and Dimitris Bertsimas. Finding neurons in a haystack: Case studies with sparse probing. *arXiv preprint arXiv:2305.01610*, 2023.
- Jiawei He, Yu Gong, Joseph Marino, Greg Mori, and Andreas Lehrmann. Variational autoencoders with jointly optimized latent dependency structure. In *International Conference on Learning Representations*, 2019.
- Irina Higgins, Loic Matthey, Arka Pal, Christopher P Burgess, Xavier Glorot, Matthew M Botvinick, Shakir Mohamed, and Alexander Lerchner. beta-VAE: Learning basic visual concepts with a constrained variational framework. *ICLR*, 3, 2017.
- Aapo Hyvärinen and Petteri Pajunen. Nonlinear independent component analysis: Existence and uniqueness results. *Neural networks*, 12(3):429–439, 1999.
- Haque Ishfaq, Assaf Hoogi, and Daniel Rubin. TVAE: Triplet-based variational autoencoder using metric learning. *arXiv preprint arXiv:1802.04403*, 2018.
- Yibo Jiang, Goutham Rajendran, Pradeep Ravikumar, Bryon Aragam, and Victor Veitch. On the origins of linear representations in large language models. *arXiv preprint*, 2024.
- Matthew J Johnson, David K Duvenaud, Alex Wiltschko, Ryan P Adams, and Sandeep R Datta. Composing graphical models with neural networks for structured representations and fast inference. *Advances in neural information processing systems*, 29, 2016.
- David Kaltenpoth and Jilles Vreeken. Nonlinear causal discovery with latent confounders. In Andreas Krause, Emma Brunskill, Kyunghyun Cho, Barbara Engelhardt, Sivan Sabato, and Jonathan Scarlett, editors, *Proceedings of the 40th International Conference on Machine Learning*, volume 202 of *Proceedings of Machine Learning Research*, pages 15639–15654. PMLR, 23–29 Jul 2023.
- Been Kim, Martin Wattenberg, Justin Gilmer, Carrie Cai, James Wexler, Fernanda Viegas, et al. Interpretability beyond feature attribution: Quantitative testing with concept activation vectors (TCAV). In *International conference on machine learning*, pages 2668–2677. PMLR, 2018.
- Hyunjik Kim and Andriy Mnih. Disentangling by factorising. In *International conference on machine learning*, pages 2649–2658. PMLR, 2018.
- Diederik P. Kingma and Max Welling. Auto-encoding variational bayes. In Yoshua Bengio and Yann LeCun, editors, *2nd International Conference on Learning Representations, ICLR 2014, Banff, AB, Canada, April 14-16, 2014, Conference Track Proceedings*, 2014. URL <http://arxiv.org/abs/1312.6114>.

- Bohdan Kivva, Goutham Rajendran, Pradeep Ravikumar, and Bryon Aragam. Identifiability of deep generative models without auxiliary information. *Advances in Neural Information Processing Systems*, 35:15687–15701, 2022.
- Murat Kocaoglu, Christopher Snyder, Alexandros G Dimakis, and Sriram Vishwanath. CausalGAN: Learning causal implicit generative models with adversarial training. *arXiv:1709.02023 [cs.LG]*, 2017.
- Sébastien Lachapelle, Pau Rodríguez, Yash Sharma, Katie Everett, Rémi Le Priol, Alexandre Lacoste, and Simon Lacoste-Julien. Disentanglement via mechanism sparsity regularization: A new principle for nonlinear ICA. In Bernhard Schölkopf, Caroline Uhler, and Kun Zhang, editors, *1st Conference on Causal Learning and Reasoning, CLear 2022*, volume 177 of *Proceedings of Machine Learning Research*, pages 428–484. PMLR, 2022. URL <https://proceedings.mlr.press/v177/lachapelle22a.html>.
- Yann LeCun, Léon Bottou, Yoshua Bengio, and Patrick Haffner. Gradient-based learning applied to document recognition. *Proceedings of the IEEE*, 86(11):2278–2324, 1998.
- Tobias Leemann, Michael Kirchhof, Yao Rong, Enkelejda Kasneci, and Gjergji Kasneci. When are post-hoc conceptual explanations identifiable? In *Uncertainty in Artificial Intelligence*, pages 1207–1218. PMLR, 2023.
- Omer Levy and Yoav Goldberg. Neural word embedding as implicit matrix factorization. *Advances in neural information processing systems*, 27, 2014.
- Kenneth Li, Oam Patel, Fernanda Viégas, Hanspeter Pfister, and Martin Wattenberg. Inference-time intervention: Eliciting truthful answers from a language model. *arXiv preprint arXiv:2306.03341*, 2023.
- Xiaopeng Li, Zhouong Chen, Leonard KM Poon, and Nevin L Zhang. Learning latent superstructures in variational autoencoders for deep multidimensional clustering. *arXiv preprint arXiv:1803.05206*, 2018.
- Xiu-Chuan Li, Kun Zhang, and Tongliang Liu. Causal structure recovery with latent variables under milder distributional and graphical assumptions. In *The Twelfth International Conference on Learning Representations*, 2024.
- Yuhang Liu, Zhen Zhang, Dong Gong, Mingming Gong, Biwei Huang, Anton van den Hengel, Kun Zhang, and Javen Qinfeng Shi. Identifiable latent polynomial causal models through the lens of change. *arXiv preprint arXiv:2310.15580*, 2023.
- Francesco Locatello, Stefan Bauer, Mario Lucic, Gunnar Raetsch, Sylvain Gelly, Bernhard Schölkopf, and Olivier Bachem. Challenging common assumptions in the unsupervised learning of disentangled representations. In *International Conference on Machine Learning*, pages 4114–4124. PMLR, 2019.
- Francesco Locatello, Ben Poole, Gunnar Rätsch, Bernhard Schölkopf, Olivier Bachem, and Michael Tschannen. Weakly-supervised disentanglement without compromises. In *International conference on machine learning*, pages 6348–6359. PMLR, 2020.

- Sarah Mameche, David Kaltenpoth, and Jilles Vreeken. Learning causal models under independent changes. *Advances in Neural Information Processing Systems*, 36:75595–75622, 2023.
- Alex Markham and Moritz Grosse-Wentrup. Measurement dependence inducing latent causal models. In Jonas Peters and David Sontag, editors, *Proceedings of the 36th Conference on Uncertainty in Artificial Intelligence (UAI)*, volume 124 of *Proceedings of Machine Learning Research*, pages 590–599. PMLR, 03–06 Aug 2020.
- Alex Markham, Mingyu Liu, Bryon Aragam, and Liam Solus. Neuro-causal factor analysis. *arXiv preprint arXiv:2305.19802*, 2023.
- Alex Markham, Isaac Hirsch, Jeri A Chang, Liam Solus, and Bryon Aragam. Intervening to learn and compose causally disentangled representations. In Bijan Mazaheri and Niels Richard Hansen, editors, *Proceedings of the Fifth Conference on Causal Learning and Reasoning*, Proceedings of Machine Learning Research. PMLR, Apr 2026.
- Nathan Juraj Michlo. Disent - a modular disentangled representation learning framework for pytorch. Github, 2021. URL <https://github.com/nmichlo/disent>.
- Tomáš Mikolov, Wen-tau Yih, and Geoffrey Zweig. Linguistic regularities in continuous space word representations. In *Proceedings of the 2013 conference of the north american chapter of the association for computational linguistics: Human language technologies*, pages 746–751, 2013.
- Zhoutong Mo et al. Compositional generative modeling: A single model is not all you need, 2024. URL <https://arxiv.org/abs/2402.01103>.
- F. Mölder, K. P. Jablonski, B. Letcher, et al. Sustainable data analysis with Snakemake. *F1000Research*, 10:33, 2025. doi: 10.12688/f1000research.29032.3.
- Milton Montero, Jeffrey Bowers, Rui Ponte Costa, Casimir Ludwig, and Gaurav Malhotra. Lost in latent space: Examining failures of disentangled models at combinatorial generalisation. *Advances in Neural Information Processing Systems*, 35:10136–10149, 2022.
- Milton Llera Montero, Casimir JH Ludwig, Rui Ponte Costa, Gaurav Malhotra, and Jeffrey Bowers. The role of disentanglement in generalisation. In *International Conference on Learning Representations*, 2021.
- Raha Moraffah, Bahman Moraffah, Mansooreh Karami, Adrienne Raglin, and Huan Liu. Causal adversarial network for learning conditional and interventional distributions. *arXiv preprint arXiv:2008.11376*, 2020.
- Gemma Elyse Moran, Dhanya Sridhar, Yixin Wang, and David Blei. Identifiable deep generative models via sparse decoding. *Transactions on Machine Learning Research*, 2022.
- Hiroshi Morioka and Aapo Hyvärinen. Causal representation learning made identifiable by grouping of observational variables. *arXiv preprint arXiv:2310.15709*, 2023.
- Luca Moschella, Valentino Maiorca, Marco Fumero, Antonio Norelli, Francesco Locatello, and Emanuele Rodola. Relative representations enable zero-shot latent space communication. *arXiv preprint arXiv:2209.15430*, 2022.

- Jacobie Mouton and Rodney Stephen Kroon. Integrating Bayesian network structure into residual flows and variational autoencoders. *Transactions on Machine Learning Research*, 2023.
- Neel Nanda, Andrew Lee, and Martin Wattenberg. Emergent linear representations in world models of self-supervised sequence models. *arXiv preprint arXiv:2309.00941*, 2023.
- Kiho Park, Yo Joong Choe, and Victor Veitch. The linear representation hypothesis and the geometry of large language models. *arXiv preprint arXiv:2311.03658*, 2023.
- Jeffrey Pennington, Richard Socher, and Christopher D Manning. Glove: Global vectors for word representation. In *Proceedings of the 2014 conference on empirical methods in natural language processing (EMNLP)*, pages 1532–1543, 2014.
- Alec Radford, Luke Metz, and Soumith Chintala. Unsupervised representation learning with deep convolutional generative adversarial networks. *arXiv preprint arXiv:1511.06434*, 2015.
- Maithra Raghu, Justin Gilmer, Jason Yosinski, and Jascha Sohl-Dickstein. SVCCA: Singular vector canonical correlation analysis for deep understanding and improvement. *stat*, 1050:19, 2017.
- Goutham Rajendran, Simon Buchholz, Bryon Aragam, Bernhard Schölkopf, and Pradeep Ravikumar. Learning interpretable concepts: Unifying causal representation learning and foundation models. *arXiv preprint arXiv:2402.09236*, 2024.
- Patrik Reizinger, Luigi Gresele, Jack Brady, Julius von Kügelgen, Dominik Zietlow, Bernhard Schölkopf, Georg Martius, Wieland Brendel, and Michel Besserve. Embrace the gap: Vaes perform independent mechanism analysis. *arXiv preprint arXiv:2206.02416*, 2022.
- Danilo Rezende and Shakir Mohamed. Variational inference with normalizing flows. In *International conference on machine learning*, pages 1530–1538. PMLR, 2015.
- Danilo Jimenez Rezende, Shakir Mohamed, and Daan Wierstra. Stochastic backpropagation and approximate inference in deep generative models. In *International conference on machine learning*, pages 1278–1286. PMLR, 2014.
- Bernhard Schölkopf, Francesco Locatello, Stefan Bauer, Nan Rosemary Ke, Nal Kalchbrenner, Anirudh Goyal, and Yoshua Bengio. Toward causal representation learning. *Proceedings of the IEEE*, 109(5):612–634, 2021.
- Lukas Schott, Julius von Kügelgen, Frederik Träuble, Peter Gehler, Chris Russell, Matthias Bethge, Bernhard Schölkopf, Francesco Locatello, and Wieland Brendel. Visual representation learning does not generalize strongly within the same domain. In *10th International Conference on Learning Representations (ICLR 2022)*, 2022.
- Yeon Seonwoo, Sungjoon Park, Dongkwan Kim, and Alice Oh. Additive compositionality of word vectors. In *Proceedings of the 5th Workshop on Noisy User-generated Text (W-NUT 2019)*, pages 387–396, 2019.
- Xinwei Shen, Furui Liu, Hanze Dong, Qing Lian, Zhitang Chen, and Tong Zhang. Weakly supervised disentangled generative causal representation learning. *Journal of Machine Learning Research*, 23:1–55, 2022.

- Patrice Y Simard, David Steinkraus, John C Platt, et al. Best practices for convolutional neural networks applied to visual document analysis. In *Icdar*, volume 3. Edinburgh, 2003.
- Casper Kaae Sønderby, Tapani Raiko, Lars Maaløe, Søren Kaae Sønderby, and Ole Winther. Ladder variational autoencoders. *Advances in neural information processing systems*, 29, 2016.
- Chandler Squires, Anna Seigal, Salil Bhate, and Caroline Uhler. Linear causal disentanglement via interventions, 2023.
- Nils Sturma, Chandler Squires, Mathias Drton, and Caroline Uhler. Unpaired multi-domain causal representation learning. *Advances in Neural Information Processing Systems*, 36:34465–34492, 2023.
- Christian Szegedy, Wojciech Zaremba, Ilya Sutskever, Joan Bruna, Dumitru Erhan, Ian Goodfellow, and Rob Fergus. Intriguing properties of neural networks. *arXiv preprint arXiv:1312.6199*, 2013.
- Davide Talon, Phillip Lippe, Stuart James, Alessio Del Bue, and Sara Magliacane. Towards the reusability and compositionality of causal representations. *arXiv preprint arXiv:2403.09830*, 2024.
- Ian Tenney, Dipanjan Das, and Ellie Pavlick. BERT rediscovers the classical NLP pipeline. *arXiv preprint arXiv:1905.05950*, 2019.
- Valentin Thomas, Emmanuel Bengio, William Fedus, Jules Pondard, Philippe Beaudoin, Hugo Larochelle, Joelle Pineau, Doina Precup, and Yoshua Bengio. Disentangling the independently controllable factors of variation by interacting with the world. In *Workshop on Learning Disentangled Representations at NeurIPS*, 2017.
- Curt Tigges, Oskar John Hollinsworth, Atticus Geiger, and Neel Nanda. Linear representations of sentiment in large language models. *arXiv preprint arXiv:2310.15154*, 2023.
- Matthew Trager, Pramuditha Perera, Luca Zancato, Alessandro Achille, Parminder Bhatia, and Stefano Soatto. Linear spaces of meanings: Compositional structures in vision-language models. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 15395–15404, 2023.
- Arash Vahdat and Jan Kautz. NVAE: A deep hierarchical variational autoencoder. *Advances in neural information processing systems*, 33:19667–19679, 2020.
- Burak Varici, Emre Acartürk, Karthikeyan Shanmugam, Abhishek Kumar, and Ali Tajer. Score-based causal representation learning with interventions. *arXiv preprint arXiv:2301.08230*, 2023.
- Burak Varici, Emre Acartürk, Karthikeyan Shanmugam, and Ali Tajer. General identifiability and achievability for causal representation learning. In *International Conference on Artificial Intelligence and Statistics*, pages 2314–2322. PMLR, 2024.
- Julius von Kügelgen, Michel Besserve, Wendong Liang, Luigi Gresele, Armin Kekić, Elias Bareinboim, David M Blei, and Bernhard Schölkopf. Nonparametric identifiability of causal representations from unknown interventions. In *Advances in Neural Information Processing Systems*, 2023.

- Yixin Wang, David Blei, and John P Cunningham. Posterior collapse and latent variable non-identifiability. *Advances in Neural Information Processing Systems*, 34, 2021.
- Zihao Wang, Lin Gui, Jeffrey Negrea, and Victor Veitch. Concept algebra for score-based conditional model. In *ICML 2023 Workshop on Structured Probabilistic Inference & Generative Modeling*, 2023.
- Stefan Webb, Adam Golinski, Rob Zinkov, Tom Rainforth, Yee Whye Teh, and Frank Wood. Faithful inversion of generative models for effective amortized inference. *Advances in Neural Information Processing Systems*, 31, 2018.
- Antoine Wehenkel and Gilles Louppe. Graphical normalizing flows. In *International Conference on Artificial Intelligence and Statistics*, pages 37–45. PMLR, 2021.
- Christian Weilbach, Boyan Beronov, Frank Wood, and William Harvey. Structured conditional continuous normalizing flows for efficient amortized inference in graphical models. In *International Conference on Artificial Intelligence and Statistics*, pages 4441–4451. PMLR, 2020.
- Matthew Willetts and Brooks Paige. I don’t need $\mathbf{u}$: Identifiable non-linear ICA without side information. *arXiv preprint arXiv:2106.05238*, 2021.
- Danru Xu, Dingling Yao, Sébastien Lachapelle, Perouz Taslakian, Julius Von Kügelgen, Francesco Locatello, and Sara Magliacane. A sparsity principle for partially observable causal representation learning. *arXiv preprint arXiv:2403.08335*, 2024.
- Zhenlin Xu, Marc Niethammer, and Colin Raffel. Compositional generalization in unsupervised compositional representation learning: A study on disentanglement and emergent language, 2022. URL <https://arxiv.org/abs/2210.00482>.
- Mengyue Yang, Furui Liu, Zhitang Chen, Xinwei Shen, Jianye Hao, and Jun Wang. CausalVAE: Disentangled representation learning via neural structural causal models. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 9593–9602, 2021.
- Dingling Yao, Danru Xu, Sébastien Lachapelle, Sara Magliacane, Perouz Taslakian, Georg Martius, Julius von Kügelgen, and Francesco Locatello. Multi-view causal representation learning with partial observability. *arXiv preprint arXiv:2311.04056*, 2023.
- Dingling Yao, Dario Rancati, Riccardo Cadei, Marco Fumero, and Francesco Locatello. Unifying causal representation learning with the invariance principle. *arXiv preprint arXiv:2409.02772*, 2024.
- Ming Yuan and Yi Lin. Model selection and estimation in regression with grouped variables. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 68(1):49–67, 2006.
- Kun Zhang, Shaoan Xie, Ignavier Ng, and Yujia Zheng. Causal representation learning from multiple distributions: A general setting. *arXiv preprint arXiv:2402.05052*, 2024.
- Roland S Zimmermann, Yash Sharma, Steffen Schneider, Matthias Bethge, and Wieland Brendel. Contrastive learning inverts the data generating process. In *International conference on machine learning*, pages 12979–12990. PMLR, 2021.

Supplementary Material: Intervening to Learn and Compose Causally Disentangled Representations

Appendix A. Related work

Our work is closely related to several parallel lines of work on structured generative models, causal representation learning, linear representations, and OOD generalization. Below we compare and contrast our contributions against this literature.

Structured generative models To provide structure such as hierarchical, graphical, causal, and disentangled structures as well as other inductive biases in the latent space, there has been a trend towards building *structured* generative models that directly impose this structure *a priori*. Early work looked at incorporating fixed, known structure into generative models, such as autoregressive, graphical, and hierarchical structure (Germain et al., 2015; Johnson et al., 2016; Sønderby et al., 2016; Webb et al., 2018; Weilbach et al., 2020; Ding et al., 2021; Mouton and Kroon, 2023). This was later translated into known *causal* structure (Kocaoglu et al., 2017; Markham et al., 2023). When the latent structure is unknown, several techniques have been developed to learn useful (not necessarily causal) structure from data (Li et al., 2018; He et al., 2019; Wehenkel and Louppe, 2021; Kivva et al., 2022; Moran et al., 2022). More recently, based on growing interest in disentangled (Bengio, 2013) and/or causal (Schölkopf et al., 2021) representation learning, methods that learn causal structure have been developed (Markham and Grosse-Wentrup, 2020; Moraffah et al., 2020; Yang et al., 2021; Ashman et al., 2022; Shen et al., 2022; Kaltenpoth and Vreeken, 2023). In contrast to this line of work, which emphasizes hard graphical constraints, we neither learn nor impose any fixed graphical structure. Instead, we emphasize performance on concept learning and composition as downstream tasks.

Causal disentanglement and representation learning Causal representation learning (Schölkopf et al., 2021) is a rapidly developing area that involves, among other goals, two key objectives: 1) Learning disentangled latent factors along with 2) Learning latent causal structure. The former is what we refer to as *causal disentanglement* (see Definition 1 and its causal interpretation). The latter problem has seen tremendous theoretical progress in recent years (Brehmer et al., 2022; Shen et al., 2022; Lachapelle et al., 2022; Moran et al., 2022; Kivva et al., 2022; Buchholz et al., 2023; Gresele et al., 2021; Ahuja et al., 2023; von Kügelgen et al., 2023; Morioka and Hyvärinen, 2023; Liu et al., 2023; Mameche et al., 2023; Yao et al., 2023; Varici et al., 2023; Sturma et al., 2023; Xu et al., 2024; Zhang et al., 2024; Li et al., 2024; Varici et al., 2024; Bing et al., 2024; Ahuja et al., 2024; Talon et al., 2024; Yao et al., 2024). Recent work has also pushed in the direction of identifying concepts (Leemann et al., 2023; Rajendran et al., 2024; Fokkema et al., 2025). Our work draws inspiration from these lines of work, which articulate precise conditions under which latent causal discovery is possible *in principle*. By contrast, our focus is on methodological aspects of causal disentanglement *in practice*. While the former is a notoriously difficult task, the latter only needs to exploit learned causal invariances: In particular, causal disentanglement is possible without learning latent causal structure.

Linear representations Generative models are known to represent concepts linearly in embedding space (e.g. Mikolov et al., 2013; Szegedy et al., 2013; Radford et al., 2015); see also Levy and

Goldberg (2014); Arora et al. (2016); Gittens et al. (2017); Allen and Hospedales (2019); Ethayarajh et al. (2018); Seonwoo et al. (2019). This phenomenon has been well-documented over the past decade in both language models (Mikolov et al., 2013; Pennington et al., 2014; Arora et al., 2016; Conneau et al., 2018; Tenney et al., 2019; Elhage et al., 2022; Burns et al., 2022; Tigges et al., 2023; Nanda et al., 2023; Moschella et al., 2022; Li et al., 2023; Park et al., 2023; Gurnee et al., 2023; Cunningham et al., 2024; Jiang et al., 2024) and computer vision (Radford et al., 2015; Raghu et al., 2017; Bau et al., 2017; Engel et al., 2017; Kim et al., 2018; Chen et al., 2020; Wang et al., 2023; Trager et al., 2023). Our approach actively exploits this tendency by searching for concept representations as linear projections of the embeddings learned by a black-box model. In contrast to this prior work, our emphasis is on applying these empirical observations for causal disentanglement, compositional generalization, and OOD generation.

OOD generalization A growing line of work studies the OOD generalization capabilities of generative models, with the general observation being that existing methods struggle to generalize OOD (Xu et al., 2022; Mo et al., 2024; Montero et al., 2022, 2021; Schott et al., 2022). It is worth noting that most if not all of this work evaluates OOD generalization using reconstruction on held-out OOD samples, as opposed to generation. For example, a traditional VAE may be able to reconstruct held-out samples, but it is not possible to actively sample OOD. See Section 2 for more discussion.

Appendix B. Identifiability of the proposed model

In this appendix we set up and prove Theorem 4.1.

B.1. Set-up

Our setup is the following: Let $\mathbf{x} = (\mathbf{x}_1, \dots, \mathbf{x}_n)$, $\mathbf{c} = (\mathbf{c}_1, \dots, \mathbf{c}_{d_c})$ and $\mathbf{e} = (\mathbf{e}_1, \dots, \mathbf{e}_{d_e})$ be random vectors taking values in $\mathcal{X}, \mathcal{C}$ and $\mathcal{E}$, respectively. For simplicity, we assume that $\mathcal{X} \subseteq \mathbb{R}^n$, $\mathbf{e} \in \mathcal{E} \subseteq \mathbb{R}^{d_e}$, and $\mathbf{x} = f(\mathbf{e})$ for some embeddings $\mathbf{e} \in \mathcal{E}$. We assume that f is injective and differentiable. This map is allowed to be arbitrarily non-linear. We make no additional assumptions on f .

The random variables $\mathbf{c}_j$ will be referred to as “concepts”: Intuitively they represent abstract concepts such as shape, size, colour, etc. that may not be perfectly represented by any particular embedding in $\mathbf{e}$. The *linear representation hypothesis* (Section A) is an empirical observation that abstract concepts can be approximately represented as linear projections of a sufficiently flexible latent space. This is operationalized by assuming that $\mathbf{c}_j \approx C_j \mathbf{e}$ (see (A1) below). The matrix C_j is referred to as the *representation* of the concept $\mathbf{c}_j$. If $\mathbf{c}_j \in \mathbb{R}$ (i.e., $C_j \in \mathbb{R}^{1 \times d_e}$), then $\mathbf{c}_j$ is referred to as an *atom* and C_j its *atomic representation*. We will often denote atoms by $\mathbf{a}_j$ for notational clarity.

We consider $(\mathbf{x}, \mathbf{c}, \mathbf{e})$ satisfying the following:

- (A1) $\mathbf{c}_j = C_j \mathbf{e} + \eta_j$ for all $j \in [d_c]$ for some real matrices $C_1 \dots, C_{d_c}$, where $\eta_j \sim \mathcal{N}(0, \Lambda_j)$ for some diagonal matrix $\Lambda_j \succ 0$.
- (A2) $\mathbf{x} = f(\mathbf{e})$ for some injective and differentiable function f .
- (A3) There is a DAG $G = ([d_c], E)$ such that

$$\mathbf{c}_j = \sum_{k \in \text{pa}_G(j)} \alpha_{kj} \mathbf{c}_k + \epsilon_j, \quad \text{for all } j \in [d_c] \quad (5)$$

and where $\varepsilon_1, \dots, \varepsilon_{d_c}$ are mutually independent and $\mathbf{e}, \eta$ are independent of $\varepsilon = (\varepsilon_1, \dots, \varepsilon_{d_c})$.

Note that these assumptions are not necessary *in practice*—as evidenced by our experiments in Section 4 which clearly violate these assumptions—and are merely used here to provide a proof-of-concept identifiability result based on our architectural choices. The representation (5) implies, in particular, the usual factorization

$$p(\mathbf{c}_1, \dots, \mathbf{c}_{d_c}) = \prod_{j=1}^{d_c} p(\mathbf{c}_j \mid \text{pa}_G(j)),$$

which formalizes the notion of causal disentanglement as in Schölkopf et al. (2021).

Our goal is to identify the concept marginal $p(\mathbf{c})$ and concept representations $(C_1, \dots, C_{d_c})$ from concept interventions. Concept interventions are well-defined through the structural causal model defined by the DAG G and its structural equations. We consider single-node concept interventions where $\mathbf{c}_j$ is stochastically set to be centered at μ_j with variance $\Omega_j \succ 0$ (making this precise requires some set-up; see Appendix B.3.1 for details).

By identifiability, we mean the usual notion of identifiability up to permutation, shift, and scaling from the literature. Namely, for any two sets of parameters $(f, C_1, \dots, C_{d_c})$ and $(\tilde{f}, \tilde{C}_1, \dots, \tilde{C}_{d_c})$ that generate the same observed marginal $p(\mathbf{x})$, there exists permutations P_j , diagonal scaling matrices D_j , and a shift $b \in \mathbb{R}^{d_e}$, such that

$$C_j f^{-1}(\mathbf{x}) = D_j P_j \tilde{C}_j (\tilde{f}^{-1}(\mathbf{x}) + b), \quad \text{for all } j \text{ and } \mathbf{x}. \quad (6)$$

Moreover, there exists an invertible linear transformation L such that the concepts satisfy

$$C_j = P_j \tilde{C}_j L^{-1}. \quad (7)$$

This definition of identifiability, which aligns with similar notions of identifiability that have appeared previously (e.g. Squires et al., 2023; von Kügelgen et al., 2023; Rajendran et al., 2024; Fokkema et al., 2025), implies in particular that the concepts $\mathbf{c}_j$ are identifiable up to permutation, shift, and scale, and that their representations are identifiable up to a linear transformation. As a result, the concept marginal $p(\mathbf{c})$ is also identifiable (again up to permutation, shift, and scale as well). The shift ambiguity can of course be resolved by assuming the concepts are centered (zero-mean), and the scale ambiguity can be resolved by assuming concepts are normalized (e.g., unit norm).

We make the following assumptions, which are adapted from Rajendran et al. (2024) to the present setting. Note that Rajendran et al. (2024) do not consider a structural causal model over $\mathbf{c}$, and thus the notion of a concept intervention there is not well-defined.

Assumption 1 (Atomic representations) *There exists a set of atomic representations $\mathcal{A} = \{\mathbf{a}_1, \dots, \mathbf{a}_n\}$ of linearly independent vectors $\mathbf{a}_j \in \mathbb{R}^{d_e}$ such that the rows of each representation C_j are in $\mathcal{A}$. We denote the indices of $\mathcal{A}$ that appear as rows of C_j by S^j , i.e., $S^j = \{i : \mathbf{a}_i \text{ is a row in } C_j\}$. We assume that $\cup_j S^j = \mathcal{A}$, i.e., every atom in $\mathcal{A}$ appears in some concept representation C_j .*

Define a matrix $M \in \mathbb{R}^{n \times d_c}$ to track which concepts use which atoms, i.e.

$$M_{ij} = \begin{cases} \frac{1}{\omega_j^2} & \text{if } i \in S^j \\ 0 & \text{otherwise,} \end{cases} \quad (8)$$

where ω_j^2 are the diagonal entries of Ω_j . Similarly, we define a matrix $B \in \mathbb{R}^{n \times d_c}$ given by

$$B_{ij} = \begin{cases} \frac{(\mu_j)_k}{\omega_j^2} & \text{if } i \in S^j \text{ and the } k\text{th row of } C_j \text{ is } \mathbf{a}_i, \\ 0 & \text{otherwise.} \end{cases} \quad (9)$$

The second assumption ensures that the concepts and associated interventional environments are sufficiently diverse.

Assumption 2 (Concept diversity) *For every pair of atoms $\mathbf{a}_i$ and $\mathbf{a}_j$ with $i \neq j$ there is a concept $\mathbf{c}_k$ such that $i \in S^k$ and $j \notin S^k$. Furthermore, $\text{rank}(M) = n$ and there exists $v \in \mathbb{R}^{d_c}$ such that $v^\top M = 0$ and $(v^\top B)_i \neq 0$ for each coordinate i .*

B.2. Formal statement of Theorem 4.1

With this setup, we can now state and prove a more formal version of Theorem 4.1.

Theorem 5 *Under Assumptions 1-2 and given single-node interventions on each concept $\mathbf{c}_j$, we can identify the representations C_j and the latent concept distribution $p(\mathbf{c})$, in the sense of (6-7).*

In the context of learning disentangled representations, identifying the concept marginal $p(\mathbf{c})$ is particularly relevant, since this implies we also learn certain (in)dependence relationships between the concepts. In other words, if the “true” concept representations are disentangled (i.e., independent), then we identify disentangled concepts.

In relation to existing work, we make the following remarks:

- While our proof ultimately relies on techniques from [Rajendran et al. \(2024\)](#), Theorem 5 does not trivially follow from their results, which do not cover interventions over the concept marginal $p(\mathbf{c})$. Extending their results to the case of concept interventions, including making sure these interventions are well-defined causal objects, is the main technical difficulty in our proof. This is surprisingly tricky since we do not assume a full structural causal model over $(\mathbf{x}, \mathbf{c}, \mathbf{e})$: Interventions are only well-defined over $\mathbf{c}$, which makes deducing the downstream effects on the observed $\mathbf{x}$ somewhat subtle.
- Various results in CRL (e.g. [Buchholz et al., 2023](#); [von Kügelgen et al., 2023](#); [Varici et al., 2024](#), see Section A for more references) prove identifiability results under interventions, however, these results require interventions directly on the embeddings $\mathbf{e}$, as opposed to the concepts $\mathbf{c}$. Theorem 5, by contrast, requires only $d_c \ll d_e$ total interventions and in doing so directly addresses the technical challenges associated with intervening directly on concepts as opposed to embeddings.
- [Leemann et al. \(2023\)](#); [Fokkema et al. \(2025\)](#) study concept identifiability when the nonlinear function f is known. We avoid this simplifying assumption (sometimes called *post-hoc* identifiability), and dealing with the potential nonidentifiability of f is another technical complication our setting must deal with. Indeed, in our setting we can only identify the behaviour of f on the concept subspaces defined by C_j , which is much weaker than identifying all of f on $\mathcal{E}$. On the other hand, we have left finite-sample aspects to future work.

B.3. Proof of Theorem 4.1

To prove Theorem 4.1, we begin with some notation. Let $\mathcal{M}$ denote the collection of all distributions P over $(\mathbf{x}, \mathbf{c}, \mathbf{e})$ satisfying (A1)-(A3) for some choice of $C_1, \dots, C_{d_c}, \eta, G, f_1, \dots, f_{d_e}$, and ε . For a fixed G , let

$$\mathcal{M}(G) := \{P = P(\mathbf{x}, \mathbf{c}, \mathbf{e}) \in \mathcal{M} : P \text{ satisfies (A3) with } G\}. \quad (10)$$

The model $\mathcal{M}(G)$ is nonempty for any G , as can be seen by letting $\mathbf{e}, \eta$ and ε be normally distributed and taking all functions to be linear. Moreover, $\mathcal{M} = \cup_G \mathcal{M}(G)$. Finally, let $\mathcal{M}_c(G)$ denote the collection of all marginal distributions over $\mathbf{c}$ induced by the distributions in $\mathcal{M}(G)$.

B.3.1. CONCEPT INTERVENTIONS

In practice, the edge structure of the graph G is unknown, and we do not assume this is known in advance. We work in the setting where we can only observe outcomes of the variables $\mathbf{x}$. However, we allow for the possibility to intervene on the variables $\mathbf{c}_1, \dots, \mathbf{c}_{d_c}$. We consider only interventional distributions with single-node targets $I = \{j\}$ for $j \in [m]$. Note that, by definition of $\mathcal{M}_c(G)$, if $P \in \mathcal{M}_c(G)$ then there is a distribution $\tilde{P} \in \mathcal{M}(G)$ such that $\tilde{p}(\mathbf{c}) = p(\mathbf{c})$ for all $\mathbf{c} \in \mathcal{C}$. Although immediate from these definitions, we record this as a lemma for future use:

Lemma 6 *If $P \in \mathcal{M}_c(G)$ then there is a distribution $\tilde{P} \in \mathcal{M}(G)$ such that $\tilde{p}(\mathbf{c}) = p(\mathbf{c})$ for all $\mathbf{c} \in \mathcal{C}$.*

For the target $I = \{j\}$ and distribution $P \in \mathcal{M}_c(G)$ we say that P^j is an *interventional distribution* for j and P if in P^j we have the following identities:

$$\begin{aligned} \mathbf{c}_i &= \sum_{k \in \text{pa}_G(i)} \alpha_{ki} \mathbf{c}_k + \varepsilon_i, \\ \mathbf{c}_j &= \boldsymbol{\mu}_j + \boldsymbol{\eta}'_j, \quad \text{where } \boldsymbol{\eta}'_j \sim \mathcal{N}(0, \omega_j^2). \end{aligned} \quad (11)$$

That is, the structural equations defining $\mathbf{c}_i$ for $i \neq j$ are equal to those defining $\mathbf{c}_i$ in the observational distribution P , while the distribution of $\mathbf{c}_j$ is manipulated to be $\mathcal{N}(\boldsymbol{\mu}_j, \omega_j^2)$, independent of the other concepts.

The following is obvious from our assumptions:

Lemma 7 *For any distribution $P \in \mathcal{M}_c(G)$, the interventional distribution P^j defined by (11) is also an element of $\mathcal{M}_c(G)$, i.e., $P^j \in \mathcal{M}_c(G)$.*

Thus, we have the usual, well-defined notion of intervention over the concepts $\mathbf{c}$. The next step is to translate the effect of these interventions onto $\mathbf{e}$ and $\mathbf{x}$.

The following proposition shows that an interventional distribution $P^j \in \mathcal{M}_c(G)$ for $P \in \mathcal{M}_c(G)$ always extends to a well-defined post-interventional joint distribution over $(\mathbf{x}, \mathbf{c}, \mathbf{e})$ that can be interpreted as an interventional distribution for the appropriate choice of joint observational distribution over $(\mathbf{x}, \mathbf{c}, \mathbf{e})$.

Proposition 8 *Let P^j be an interventional distribution for j and $P \in \mathcal{M}_c(G)$. Then there exists a distribution $\tilde{P}^j \in \mathcal{M}(G)$ such that*

1. $\tilde{p}^j(\mathbf{c}) = p^j(\mathbf{c})$ for all $\mathbf{c} \in \mathcal{C}$,

2. $\tilde{p}^j(\mathbf{e}) = \tilde{p}(\mathbf{e})$ for all $\mathbf{e} \in \mathcal{E}$,
3. If $\tilde{P}$ is induced by the functional relation $\mathbf{x} = f(\mathbf{e})$, and similarly $\tilde{P}^j$ with $\mathbf{x} = f^j(\mathbf{e})$ then $f^j = f$.

Proof Lemma 7 implies that $P^j \in \mathcal{M}_{\mathbf{c}}(\mathcal{G})$. Thus Lemma 6 implies the existence of joint distributions $\tilde{P}, \tilde{P}^j \in \mathcal{M}(\mathcal{G})$ such that

$$\tilde{p}(\mathbf{c}) = p(\mathbf{c}) \quad \text{and} \quad \tilde{p}^j(\mathbf{c}) = p^j(\mathbf{c}) \text{ for all } \mathbf{c} \in \mathcal{C}. \quad (12)$$

(1) follows from the definition of $\tilde{P}^j \in \mathcal{M}(\mathcal{G})$ for $P^j \in \mathcal{M}_{\mathbf{c}}(\mathcal{G})$ from Lemma 6 above; i.e., (12) above.

(2) follows from the definition of an interventional distribution and the factorization

$$\tilde{p}(\mathbf{x}, \mathbf{c}, \mathbf{e}) = \tilde{p}(\mathbf{x} \mid \mathbf{e})\tilde{p}(\mathbf{c} \mid \mathbf{e})\tilde{p}(\mathbf{e}),$$

satisfied by any distribution $P \in \mathcal{M}(\mathcal{G})$. Namely, the interventional distribution P^j satisfies

$$\begin{aligned} \tilde{p}^j(\mathbf{x}, \mathbf{c}, \mathbf{e}) &= \tilde{p}^j(\mathbf{x} \mid \mathbf{e})\tilde{p}^j(\mathbf{c} \mid \mathbf{e})\tilde{p}^j(\mathbf{e}), \\ &= \tilde{p}^j(\mathbf{x} \mid \mathbf{e})p^j(\mathbf{c} \mid \mathbf{e})\tilde{p}^j(\mathbf{e}), \end{aligned}$$

according to the definition of $\tilde{P}$. Since the intervention only perturbs the conditional factors $\tilde{p}(\mathbf{c}_j \mid \mathbf{c}_{\text{pa}_{\mathcal{G}}(j)}, \mathbf{e})$ in $\tilde{p}(\mathbf{c} \mid \mathbf{e})$ for j , we can always choose $\tilde{P}$ and $\tilde{P}^j$ such that $\tilde{p}^j(\mathbf{e}) = \tilde{p}(\mathbf{e})$.

(3) follows similarly to (2). Namely, since the intervention only perturbs the conditional factors specified above, we can always choose $f^j = f$. ■

We additionally have the following:

Proposition 9 *Let $I_j := \{j\}$ be the collection of single-node intervention targets, and consider the interventional distributions $P^1, \dots, P^{d_{\mathbf{c}}}$ for $P \in \mathcal{M}_{\mathbf{c}}(\mathcal{G})$. Then there exists $\tilde{P}, \tilde{P}^1, \dots, \tilde{P}^{d_{\mathbf{c}}} \in \mathcal{M}(\mathcal{G})$ such that*

$$\tilde{p}^j(\mathbf{e}) = \tilde{p}(\mathbf{e}) \quad \text{for all } j \in [d_{\mathbf{c}}]$$

and $f^j = f$ for all j .

Proof Fix a choice of $\tilde{P} \in \mathcal{M}(\mathcal{G})$ for P , which specifies a function f . For the chosen $\tilde{P}$, we construct $\tilde{P}, \tilde{P}^1, \dots, \tilde{P}^{d_{\mathbf{c}}}$ according to Proposition 8. The result follows by applying Proposition 8 for every j . ■

Proposition 9 will allow us (below) to define more formally the environments that are used to identify concepts. Intuitively, each environment corresponds to single-node interventions on a single concept $\mathbf{c}_j$, however, we can only observe the effect this intervention has on the observed $\mathbf{x}$. The purpose of Propositions 8-9 are to trace the dependence between $\mathbf{x}$ and $\mathbf{c}$ —via the embeddings $\mathbf{e}$ —according to the model (A1)-(A3).

Formally, let $I_j = \{j\}$ with $\eta'_j \sim \mathcal{N}(0, \omega_j^2)$ and $\mathbf{c}_j = \boldsymbol{\mu}_j + \eta'_j$ according to (11), i.e., each environment corresponds to a single-node intervention on the j th concept $\mathbf{c}_j$. By Proposition 9, there exist post-intervention joint distributions $\tilde{P}^j \in \mathcal{M}(\mathcal{G})$ over $(\mathbf{x}, \mathbf{c}, \mathbf{e})$ such that

$$\tilde{p}^j(\mathbf{e}) = \tilde{p}(\mathbf{e}) \quad \text{for all } j \in [d_{\mathbf{c}}]$$

and $f^j = f$ for all j . Then the j th environment corresponds to sampling $\mathbf{x} = f(\mathbf{e})$ where $\mathbf{e} \sim p_j(\mathbf{e}) := \tilde{p}^j(\mathbf{e} \mid \mathbf{c}_j = \boldsymbol{\mu}_j)$.

B.3.2. REDUCTION TO CONCEPT IDENTIFICATION

We begin by computing the log-likelihood ratio between the observational and interventional environments. The idea is that the concept representations are revealed through fluctuations in this likelihood ratio. Using Proposition 9, we have

$$\log p_0(\mathbf{e}) - \log p_j(\mathbf{e}) \propto \log p_0(\mathbf{e}) - \log \tilde{p}^j(\mathbf{c}_j = \boldsymbol{\mu}_j \mid \mathbf{e}) - \log \tilde{p}^j(\mathbf{e}) \quad (13)$$

$$\propto -\frac{1}{2}(\boldsymbol{\mu}_j - C_j \mathbf{e})^T \Omega_j^{-1} (\boldsymbol{\mu}_j - C_j \mathbf{e}) \quad (14)$$

$$\propto \sum_{i=1}^n \frac{1}{2} M_{ij} (\mathbf{a}_i^T \mathbf{e})^2 - B_{ij} \mathbf{a}_i^T \mathbf{e} + O(1). \quad (15)$$

From here, we proceed as in Rajendran et al. (2024) (note that in their notation, $\mathbf{e} = \mathbf{z}$). For completeness, we sketch the main steps here. The basic idea is to analyze the system of d_c equations induced by the log-likelihood ratios for each concept intervention. First, we reduce the system to a standard form by transforming the atoms into the standard basis via a change of coordinates. After this change of coordinates, the system can be rewritten as

$$h_j(\mathbf{e}) := \log p_0(\mathbf{e}) - \log p_j(\mathbf{e}) \propto \sum_{i=1}^n \frac{1}{2} M_{ij} \mathbf{e}_i^2 - B_{ij} \mathbf{e}_i + O(1). \quad (16)$$

These functions are convex, which allows us to identify the $|S_T|$ where $S_T = \cup_{j \in T} S_j$ by minimizing the convex function $\sum_{j \in T} h_j(\mathbf{e})$. Then a similar induction argument identifies the matrices M and B . Now, given two distinct representations of $p(\mathbf{x})$ via nonlinear mixing functions f and $\tilde{f}$, define $\varphi = \tilde{f}^{-1} \circ f$. Let $\mathbf{h} = (h_1, \dots, h_n)$ and write $\mathbf{e}^c = (\mathbf{e}_1, \dots, \mathbf{e}_n)$ and $\mathbf{e} = (\mathbf{e}^c, \mathbf{e}^\perp) \in \mathbb{R}^{n \times (d_c - n)}$. Let $\iota^\perp(\mathbf{e}^c) = (\mathbf{e}^c, 0)$, $\pi^c(\mathbf{e}) = \mathbf{e}^c$, and $\varphi^\perp : \mathbb{R}^n \rightarrow \mathbb{R}^n$ by $\varphi^\perp(\mathbf{e}^c)_i = \varphi(\mathbf{e}, 0)_i$. Note that via the change of variables formula, we can write $\mathbf{h}(\mathbf{e}) = \mathbf{H}(\mathbf{x})$ for some function $\mathbf{H}$. Then

$$\mathbf{h}(\mathbf{e}^c, 0) = \mathbf{H}(f(\mathbf{e}^c, 0)) = \mathbf{H}(\tilde{f}(\varphi^\perp(\mathbf{e}^c))) = \mathbf{h}(\varphi^\perp(\mathbf{e}^c)), \quad (17)$$

which is the crucial relation used in the proof of Theorem 1 of Rajendran et al. (2024). From here, identifiability follows from a straightforward but lengthy linear algebraic argument based on the first-order conditions for minimizing $\mathbf{h}$ and the identifiability of M, B proved above. The proof of Theorem 5 follows.

Appendix C. quad: A semi-synthetic benchmark for compositional generation

To provide a controllable, synthetic test bed for evaluating composition in a visual environment, we developed a simple semi-synthetic benchmark, visualized in Figures 3–4. quad is a visual environment defined by 8 concepts:

1. quad1: The colour of the first quadrant;
2. quad2: The colour of the second quadrant;
3. quad3: The colour of the third quadrant;
4. quad4: The colour of the fourth quadrant;
5. size: The size of the center object;

6. `orientation`: The orientation (angle) of the center object;
7. `object`: The colour of the center object;
8. `shape`: The shape of the center object (circle, square, pill, triangle).

With the exception of `shape`, which is discrete, the remaining concepts take values in $[0, 1]$. This environment has the following appealing properties:

- Composing multiple concepts is straightforward and perceptually distinct;
- Concepts are continuous, which allows a variety of soft and hard interventions;
- Sampling is fast and easy, and requires only a few lines of Python code and no external software;
- Creating OOD hold-out validation datasets of any size is straightforward;
- This can be simulated on any $n \times n$ grid (as long as n is large enough to distinguish different shapes; $n \geq 16$ is large enough in practice). In our experiments, we used $n = 64$.

These properties make `quad` especially useful for both quantitative and qualitative evaluation; indeed, it was created precisely for this reason. Example images are shown in Figures 3-4.

To describe the experimental setting used in the `quad` (independent) experiments, let c_j indicate the j th concept as listed above. We used seven different contexts, including an observational context and one context each for single-concept interventions on $(c_1, \dots, c_6)$, i.e., `quad1`, `quad2`, `quad3`, `quad4`, `size`, and `orientation`. We use a sample size of 50 000 images per context. In each context, c_7 (`object`) and c_8 (`shape`) are sampled uniformly at random. The remaining six concepts were sampled as follows:

- The observational context was generated by randomly sampling $(c_1, \dots, c_6)$ independently from the range $[0, 0.5]$,
- The interventional contexts were generated by isolating a particular concept and sampling it uniformly from $[0.5, 1]$, while sampling the rest from $[0, 0.5]$.

For example, colours in $[0, 0.5]$ include greens, yellows, and reds, and colours in $[0.5, 1]$ include blues, purples, and magenta. Thus, the object colour takes on any value in every context, whereas the four quadrant colours are restricted in the observational context. When we intervene on a particular quadrant, we see that the colour palette noticeably shifts from greens to blues (Figure 3).

Similarly, we can compose multiple interventions together, as in Figure 4. This facilitates evaluation of the concept composition capabilities of models trained on single interventions, providing a ground truth distribution to compare against generated OOD samples like in Figure 5.

We also made use of the easy-to-generate and perceptually distinct concept compositions in this dataset by evaluating models on held-out contexts (each containing 10 000 images) with double-concept interventions, including `quad2_quad3`, `quad2_quad4`, `quad2_size`, `quad1_quad2`, `quad1_quad3`, `quad1_quad4`, and `quad1_size`. Examples can be seen in Figure 4.

For comparison, see Figure E.1.2 for examples of generated multi-concept compositions (OOD) from a simple 3-layer convolutional model. The training data did not contain any multi-concept compositions.

For `quad` dependent, a random SEM is generated using the `sampler` package (Gamella et al., 2022), which also allows for data generation under shift interventions.

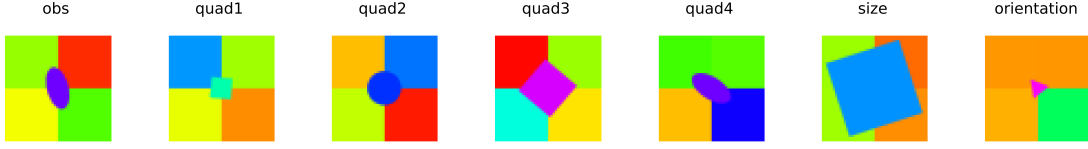


Figure 3: Example images obtained from the `quad` dataset. Each subfigure corresponds to a different context, indicated by the title (`obs` = observational; the others indicate a single-node intervention). For example, in `quad1`, the first (top left) quadrant has been manipulated from orange-green hues to blue-red hues.



Figure 4: Example images obtained from the `quad` dataset. Each subfigure corresponds to a different double-concept intervention context, indicated by the title. For example, in `quad2_quad3`, the second (top-right) and third (bottom-left) quadrants have been manipulated from orange-green hues to blue-red hues.

Appendix D. Experimental details

In this appendix we provide additional details on our experimental protocol and set-up. We have used Snakemake (Mölder et al., 2025) for reproducibility.

D.1. Evaluation metrics

We evaluated each model on the following metrics:

- **Validation ELBO loss:** This is the standard ELBO loss used for validating VAEs. We report this in bits per dimension (bpd), which is the loss scaled by a factor of $\frac{1}{\text{num_pixels} \cdot \ln 2}$.
- **p -sliced Wasserstein (SW) distance:** Unlike the loss above, which evaluates *individual reconstructed* images with respect to their corresponding ground truth images, this metric (Bonneel et al., 2015) evaluates a *generated distribution* μ_{gen} with respect to a ground truth distribution μ . Specifically, we use the implementation of Flamary et al. (2021) to compute a Monte-Carlo approximation of $SWD_p(\mu_{\text{gen}}, \mu)$, where

$$SWD_p(\mu, \nu) = \mathbb{E}_{\theta \sim \mathcal{U}(\mathbb{S}^{d-1})} \left(\mathcal{W}_p^p(\theta_{\#}\mu, \theta_{\#}\nu) \right)^{\frac{1}{p}}, \quad (18)$$

and $\theta_{\#}\mu$ stands for the pushforwards of the projection $X \in \mathbb{R}^d \mapsto \langle \theta, X \rangle$ and $p = 2$.

When reporting values for these metrics in tables, we show the mean, averaged over different runs/seeds, followed by the standard error. We generally use a 70/30 training/validation split.

D.2. Datasets

In addition to the `quad` benchmark introduced in Appendix C, we used the following benchmark datasets.

3DIdent The 3DIdent dataset (Zimmermann et al., 2021) consists of a 3D object, rendered using the Blender rendering engine, under different lighting conditions and orientations. We generated 50 000 samples per context corresponding to single-node interventions on background colour, spotlight colour, and object colour, in addition to the observational context.

MNIST We use a variation of the classic MNIST dataset (LeCun et al., 1998) known as Morpho-MNIST (Castro et al., 2019) along with additional affine transformations of MNIST (e.g., as in Simard et al. (2003)). In addition to the observational context (standard MNIST images), the dataset contains the five contexts corresponding to single-concept interventions on concepts `scaled`, `shear`, `swel`, `thic`, and `thin`. The datasets contain 60 000 images per context.

D.3. Architectures

We tested our context module end-to-end using two black-box models:

1. (lightweight) β -VAE: A standard 3-layer convolutional network with latent dimension $\dim(\mathbf{z}) = 128$. This is used for the `quad` experiments in Section 4 and the `quad` and MNIST experiments in Section E.1.
2. NVAE: A deep hierarchical VAE; see Vahdat and Kautz (2020) for details. This is used for the MNIST and 3DIdent experiments in Sections 4 and E.2.

The lightweight-VAE was used for comprehensive ablations and experiments on simple datasets, totaling more than 350 models overall, with 5 per column per table in Appendix E.1. Since NVAE requires substantially more resources and time to train, this was used to train 12 total models, shown in Table 1, the three CM variations plus three ablations on both MNIST and 3DIdent.

D.4. Additional ablations

We conducted additional ablations to stress test our model and isolate the effect of the context module. This included additional experiments on different models, regularization, and expressivity.

Models We ran ablations over different base VAEs augmented with the context module, denoted CM (end-to-end)—this is our module attached to the black-box, making full use of the different contexts, trained end-to-end. To evaluate the proposed module, we compared its performance to three natural baselines:

1. **CM pooled**: CM (end-to-end) but trained with a single context on the full *pooled* dataset—this is our module attached to the black-box given all the data but not making use of context-specific information, providing an ablation that circumvents our intervention layer (Section 3.2) while maintaining the same architectural complexity.
2. **base**: Base VAE (without the context module) trained only on a single observational context—this is the black-box when denied all (non-observational) context information and data.

3. **base pooled:** Base VAE (without the context module) trained with a single context on the full pooled dataset—this is the black-box when given all data but no explicit context information.

Results for β VAE as the base applied to `quad` and MNIST are shown in Section 4 and Appendix E.1. Results for NVAE as the base applied to 3DIdent are shown in Section 4.

Regularization We run CM (end-to-end) with β -VAE base on MNIST, independently trying out two different regularizers, varying the regularization weight λ for each: group lasso and ℓ_2 . These are applied to the weights of the reduced-form SEM A_0 in 2. We also try varying β (as in β -VAE). Results are shown in Appendix E.1.3.

Expressivity We run CM (end-to-end) on MNIST, jointly varying three parameters that control expressiveness of our context module: expressive input width w_{exp} (which controls latent dimension $\dim(\mathbf{z})$), depth of the expressive layer h_{exp} , and concept width ($w_c = \dim(\mathbf{c}_j) = w_\varepsilon = \dim(\varepsilon_j)$) as in Section 3.2). Results are shown in Appendix E.1.5.

D.5. Additional baselines

In addition to the above base models and ablations, we include comparisons against three disentanglement methods, including the unsupervised β -TCVAE (Chen et al., 2018), the weakly supervised Ada-GVAE (Locatello et al., 2020), and the supervised TVAE (Ishfaq et al., 2018). We gratefully acknowledge the open source implementation of these within the `disent` Python package (Michlo, 2021). In addition to comparisons against the `quad` ablations, these baselines are run on MNIST and 3DIdent in Appendix E.1.1.

Appendix E. Additional results

In this section, we present complete results, including experimental settings described in Appendix D for our context module with the lightweight-VAE black-box architecture as well as additional results on MNIST and 3DIdent for our context module with NVAE.

E.1. lightweight-VAE

E.1.1. ABLATIONS (MNIST AND `QUAD`)

A description of the architecture and experiment is given in Appendix D. Tables 2–5 are already summarized in the main text in Table 1, but here we use the extra space here to additionally present standard errors and show a breakdown of the results over the different contexts.

Table 2: Evaluation of the context module, ablations, and baselines in terms of reconstruction and concept learning on MNIST.

	CM (end-to-end)	CM (fine-tune)	CM (frozen)	base	base pooled	CM pooled	Ada-GVAE	BetaTCVAE	TVAE
reconstruction (ELBO BPD)									
in-distribution	0.259 $\pm$ 0.005	0.235 $\pm$ 0.004	0.231 $\pm$ 0.003	0.181 $\pm$ 0.000	0.214 $\pm$ 0.003	0.228 $\pm$ 0.002	1.035 $\pm$ 0.002	1.036 $\pm$ 0.012	1.037 $\pm$ 0.001
out-of-distribution	—	—	—	—	—	—	—	—	—
concept learning (sliced Wasserstein)									
obs	0.099 $\pm$ 0.006	0.085 $\pm$ 0.002	0.082 $\pm$ 0.001	0.067 $\pm$ 0.002	0.085 $\pm$ 0.003	0.082 $\pm$ 0.003	0.521 $\pm$ 0.009	0.502 $\pm$ 0.014	0.463 $\pm$ 0.020
scaled	0.062 $\pm$ 0.002	0.100 $\pm$ 0.007	0.109 $\pm$ 0.011	0.179 $\pm$ 0.006	0.142 $\pm$ 0.010	0.165 $\pm$ 0.008	0.526 $\pm$ 0.017	0.492 $\pm$ 0.012	0.490 $\pm$ 0.016
shear	0.119 $\pm$ 0.006	0.110 $\pm$ 0.002	0.110 $\pm$ 0.003	0.105 $\pm$ 0.002	0.109 $\pm$ 0.002	0.109 $\pm$ 0.003	0.510 $\pm$ 0.013	0.476 $\pm$ 0.012	0.488 $\pm$ 0.024
shift	0.113 $\pm$ 0.005	0.112 $\pm$ 0.003	0.105 $\pm$ 0.003	0.110 $\pm$ 0.004	0.115 $\pm$ 0.002	0.116 $\pm$ 0.004	0.494 $\pm$ 0.014	0.447 $\pm$ 0.012	0.501 $\pm$ 0.009
swel	0.126 $\pm$ 0.007	0.113 $\pm$ 0.005	0.115 $\pm$ 0.002	0.102 $\pm$ 0.003	0.129 $\pm$ 0.009	0.114 $\pm$ 0.003	0.473 $\pm$ 0.011	0.492 $\pm$ 0.007	0.481 $\pm$ 0.011
thic	0.178 $\pm$ 0.007	0.184 $\pm$ 0.002	0.193 $\pm$ 0.010	0.196 $\pm$ 0.008	0.231 $\pm$ 0.014	0.207 $\pm$ 0.014	0.474 $\pm$ 0.008	0.475 $\pm$ 0.012	0.472 $\pm$ 0.013
thin	0.064 $\pm$ 0.004	0.080 $\pm$ 0.006	0.092 $\pm$ 0.010	0.133 $\pm$ 0.011	0.107 $\pm$ 0.006	0.126 $\pm$ 0.007	0.501 $\pm$ 0.014	0.506 $\pm$ 0.013	0.504 $\pm$ 0.011
concept learning mean	0.109 $\pm$ 0.006	0.112 $\pm$ 0.004	0.115 $\pm$ 0.006	0.127 $\pm$ 0.005	0.131 $\pm$ 0.007	0.131 $\pm$ 0.006	0.500 $\pm$ 0.012	0.484 $\pm$ 0.012	0.486 $\pm$ 0.015

Table 3: Evaluation of the context module, ablations, and baselines in terms of reconstruction, concept learning, and composition on quad (independent).

	CM (end-to-end)	CM (fine-tune)	CM (frozen)	base	base pooled	CM pooled	Ada-GVAE	BetaTCVAE	TVAE
reconstruction (ELBO BPD)									
in-distribution	0.519 $\pm$ 0.022	0.518 $\pm$ 0.006	0.559 $\pm$ 0.011	0.444 $\pm$ 0.000	0.446 $\pm$ 0.002	0.840 $\pm$ 0.072	2.017 $\pm$ 0.006	1.973 $\pm$ 0.001	2.048 $\pm$ 0.008
out-of-distribution	0.720 $\pm$ 0.069	0.675 $\pm$ 0.039	0.774 $\pm$ 0.039	3.095 $\pm$ 0.295	0.487 $\pm$ 0.009	0.830 $\pm$ 0.042	2.092 $\pm$ 0.007	1.980 $\pm$ 0.003	2.115 $\pm$ 0.019
concept learning (sliced Wasserstein)									
obs	0.056 $\pm$ 0.003	0.051 $\pm$ 0.004	0.078 $\pm$ 0.005	0.061 $\pm$ 0.004	0.100 $\pm$ 0.007	0.110 $\pm$ 0.013	0.332 $\pm$ 0.007	0.317 $\pm$ 0.009	0.349 $\pm$ 0.007
quad1	0.062 $\pm$ 0.003	0.056 $\pm$ 0.005	0.075 $\pm$ 0.003	0.233 $\pm$ 0.004	0.201 $\pm$ 0.014	0.266 $\pm$ 0.020	0.342 $\pm$ 0.009	0.335 $\pm$ 0.012	0.324 $\pm$ 0.013
quad2	0.061 $\pm$ 0.005	0.058 $\pm$ 0.007	0.086 $\pm$ 0.009	0.234 $\pm$ 0.014	0.196 $\pm$ 0.012	0.254 $\pm$ 0.017	0.346 $\pm$ 0.007	0.323 $\pm$ 0.006	0.334 $\pm$ 0.010
quad3	0.063 $\pm$ 0.002	0.064 $\pm$ 0.005	0.085 $\pm$ 0.006	0.251 $\pm$ 0.013	0.210 $\pm$ 0.008	0.244 $\pm$ 0.020	0.333 $\pm$ 0.013	0.343 $\pm$ 0.010	0.334 $\pm$ 0.011
quad4	0.065 $\pm$ 0.007	0.057 $\pm$ 0.004	0.100 $\pm$ 0.011	0.227 $\pm$ 0.010	0.227 $\pm$ 0.005	0.246 $\pm$ 0.011	0.347 $\pm$ 0.016	0.334 $\pm$ 0.011	0.324 $\pm$ 0.016
size	0.075 $\pm$ 0.007	0.073 $\pm$ 0.006	0.101 $\pm$ 0.008	0.157 $\pm$ 0.009	0.122 $\pm$ 0.001	0.177 $\pm$ 0.020	0.327 $\pm$ 0.004	0.307 $\pm$ 0.010	0.336 $\pm$ 0.009
orientation	0.057 $\pm$ 0.003	0.053 $\pm$ 0.003	0.077 $\pm$ 0.006	0.060 $\pm$ 0.003	0.104 $\pm$ 0.009	0.114 $\pm$ 0.016	0.316 $\pm$ 0.011	0.349 $\pm$ 0.013	0.334 $\pm$ 0.010
concept learning mean	0.063 $\pm$ 0.004	0.059 $\pm$ 0.005	0.086 $\pm$ 0.007	0.175 $\pm$ 0.008	0.166 $\pm$ 0.008	0.202 $\pm$ 0.017	0.335 $\pm$ 0.010	0.330 $\pm$ 0.010	0.334 $\pm$ 0.011
composition (sliced Wasserstein)									
(quad1, quad2)	0.100 $\pm$ 0.028	0.108 $\pm$ 0.016	0.137 $\pm$ 0.014	0.311 $\pm$ 0.009	0.276 $\pm$ 0.014	0.343 $\pm$ 0.018	0.346 $\pm$ 0.012	0.346 $\pm$ 0.004	0.334 $\pm$ 0.007
(quad1, quad3)	0.078 $\pm$ 0.016	0.107 $\pm$ 0.011	0.146 $\pm$ 0.013	0.323 $\pm$ 0.011	0.273 $\pm$ 0.007	0.349 $\pm$ 0.008	0.367 $\pm$ 0.016	0.372 $\pm$ 0.009	0.337 $\pm$ 0.008
(quad1, quad4)	0.094 $\pm$ 0.032	0.091 $\pm$ 0.010	0.179 $\pm$ 0.014	0.346 $\pm$ 0.017	0.277 $\pm$ 0.017	0.325 $\pm$ 0.023	0.377 $\pm$ 0.018	0.343 $\pm$ 0.012	0.346 $\pm$ 0.010
(quad1, size)	0.101 $\pm$ 0.010	0.087 $\pm$ 0.005	0.111 $\pm$ 0.010	0.266 $\pm$ 0.012	0.226 $\pm$ 0.013	0.284 $\pm$ 0.011	0.343 $\pm$ 0.005	0.331 $\pm$ 0.009	0.336 $\pm$ 0.009
(quad1, orientation)	0.060 $\pm$ 0.003	0.059 $\pm$ 0.004	0.080 $\pm$ 0.002	0.236 $\pm$ 0.012	0.194 $\pm$ 0.017	0.240 $\pm$ 0.004	0.367 $\pm$ 0.016	0.360 $\pm$ 0.007	0.356 $\pm$ 0.007
(quad2, quad3)	0.067 $\pm$ 0.007	0.098 $\pm$ 0.017	0.148 $\pm$ 0.014	0.338 $\pm$ 0.017	0.283 $\pm$ 0.014	0.325 $\pm$ 0.016	0.350 $\pm$ 0.010	0.349 $\pm$ 0.012	0.350 $\pm$ 0.008
(quad2, quad4)	0.068 $\pm$ 0.004	0.078 $\pm$ 0.010	0.194 $\pm$ 0.018	0.302 $\pm$ 0.013	0.291 $\pm$ 0.015	0.332 $\pm$ 0.012	0.382 $\pm$ 0.007	0.342 $\pm$ 0.005	0.340 $\pm$ 0.013
(quad2, size)	0.099 $\pm$ 0.008	0.083 $\pm$ 0.008	0.102 $\pm$ 0.007	0.248 $\pm$ 0.008	0.197 $\pm$ 0.004	0.281 $\pm$ 0.021	0.334 $\pm$ 0.007	0.328 $\pm$ 0.015	0.336 $\pm$ 0.006
(quad2, orientation)	0.063 $\pm$ 0.004	0.058 $\pm$ 0.003	0.090 $\pm$ 0.012	0.225 $\pm$ 0.011	0.214 $\pm$ 0.010	0.274 $\pm$ 0.013	0.329 $\pm$ 0.011	0.348 $\pm$ 0.010	0.333 $\pm$ 0.008
composition mean	0.081 $\pm$ 0.013	0.086 $\pm$ 0.009	0.132 $\pm$ 0.012	0.288 $\pm$ 0.012	0.248 $\pm$ 0.012	0.306 $\pm$ 0.014	0.355 $\pm$ 0.011	0.346 $\pm$ 0.009	0.341 $\pm$ 0.008

Table 4: Evaluation of the context module, ablations, and baselines in terms of reconstruction, concept learning, and composition on quad (dependent).

	CM (end-to-end)	CM (fine-tune)	CM (frozen)	base	base pooled	CM pooled	Ada-GVAE	BetaTCVAE	TVAE
reconstruction (ELBO BPD)									
in-distribution	0.828 $\pm$ 0.004	0.840 $\pm$ 0.005	0.889 $\pm$ 0.017	0.767 $\pm$ 0.001	0.746 $\pm$ 0.000	0.752 $\pm$ 0.000	2.443 $\pm$ 0.002	2.408 $\pm$ 0.000	2.455 $\pm$ 0.002
out-of-distribution	0.863 $\pm$ 0.006	0.908 $\pm$ 0.005	0.959 $\pm$ 0.011	0.815 $\pm$ 0.005	0.746 $\pm$ 0.001	0.753 $\pm$ 0.001	2.450 $\pm$ 0.002	2.406 $\pm$ 0.000	2.463 $\pm$ 0.002
concept learning (sliced Wasserstein)									
obs	0.107 $\pm$ 0.005	0.103 $\pm$ 0.008	0.116 $\pm$ 0.003	0.110 $\pm$ 0.005	0.104 $\pm$ 0.006	0.113 $\pm$ 0.004	0.365 $\pm$ 0.007	0.375 $\pm$ 0.008	0.357 $\pm$ 0.005
quad1	0.100 $\pm$ 0.004	0.102 $\pm$ 0.004	0.107 $\pm$ 0.008	0.158 $\pm$ 0.003	0.163 $\pm$ 0.004	0.137 $\pm$ 0.007	0.373 $\pm$ 0.004	0.352 $\pm$ 0.004	0.374 $\pm$ 0.008
quad2	0.103 $\pm$ 0.004	0.100 $\pm$ 0.006	0.110 $\pm$ 0.009	0.179 $\pm$ 0.007	0.165 $\pm$ 0.008	0.173 $\pm$ 0.009	0.373 $\pm$ 0.004	0.359 $\pm$ 0.003	0.366 $\pm$ 0.005
quad3	0.127 $\pm$ 0.017	0.127 $\pm$ 0.008	0.139 $\pm$ 0.008	0.200 $\pm$ 0.010	0.214 $\pm$ 0.009	0.213 $\pm$ 0.015	0.379 $\pm$ 0.005	0.374 $\pm$ 0.006	0.372 $\pm$ 0.009
quad4	0.145 $\pm$ 0.010	0.140 $\pm$ 0.007	0.131 $\pm$ 0.011	0.305 $\pm$ 0.012	0.277 $\pm$ 0.010	0.278 $\pm$ 0.015	0.401 $\pm$ 0.015	0.386 $\pm$ 0.008	0.388 $\pm$ 0.013
size	0.109 $\pm$ 0.006	0.105 $\pm$ 0.006	0.117 $\pm$ 0.004	0.116 $\pm$ 0.004	0.121 $\pm$ 0.007	0.119 $\pm$ 0.004	0.361 $\pm$ 0.006	0.362 $\pm$ 0.005	0.367 $\pm$ 0.005
orientation	0.123 $\pm$ 0.003	0.109 $\pm$ 0.007	0.124 $\pm$ 0.004	0.120 $\pm$ 0.009	0.123 $\pm$ 0.008	0.130 $\pm$ 0.007	0.369 $\pm$ 0.006	0.353 $\pm$ 0.002	0.365 $\pm$ 0.007
concept learning mean	0.116 $\pm$ 0.007	0.112 $\pm$ 0.007	0.121 $\pm$ 0.007	0.170 $\pm$ 0.007	0.167 $\pm$ 0.007	0.166 $\pm$ 0.009	0.375 $\pm$ 0.007	0.366 $\pm$ 0.005	0.370 $\pm$ 0.007
composition (sliced Wasserstein)									
(quad1, quad2)	0.111 $\pm$ 0.008	0.097 $\pm$ 0.007	0.123 $\pm$ 0.011	0.229 $\pm$ 0.003	0.214 $\pm$ 0.005	0.196 $\pm$ 0.013	0.366 $\pm$ 0.006	0.381 $\pm$ 0.010	0.367 $\pm$ 0.007
(quad1, quad3)	0.134 $\pm$ 0.008	0.117 $\pm$ 0.003	0.140 $\pm$ 0.008	0.227 $\pm$ 0.009	0.228 $\pm$ 0.008	0.198 $\pm$ 0.011	0.390 $\pm$ 0.009	0.367 $\pm$ 0.011	0.374 $\pm$ 0.004
(quad1, quad4)	0.133 $\pm$ 0.007	0.163 $\pm$ 0.010	0.163 $\pm$ 0.008	0.284 $\pm$ 0.015	0.292 $\pm$ 0.006	0.276 $\pm$ 0.013	0.403 $\pm$ 0.012	0.376 $\pm$ 0.013	0.361 $\pm$ 0.006
(quad1, size)	0.108 $\pm$ 0.003	0.106 $\pm$ 0.004	0.123 $\pm$ 0.007	0.172 $\pm$ 0.009	0.170 $\pm$ 0.009	0.157 $\pm$ 0.011	0.360 $\pm$ 0.008	0.357 $\pm$ 0.008	0.367 $\pm$ 0.003
(quad1, orientation)	0.110 $\pm$ 0.007	0.108 $\pm$ 0.007	0.125 $\pm$ 0.006	0.167 $\pm$ 0.006	0.157 $\pm$ 0.007	0.143 $\pm$ 0.007	0.365 $\pm$ 0.010	0.371 $\pm$ 0.007	0.369 $\pm$ 0.004
(quad2, quad3)	0.158 $\pm$ 0.007	0.150 $\pm$ 0.013	0.164 $\pm$ 0.005	0.222 $\pm$ 0.008	0.200 $\pm$ 0.011	0.195 $\pm$ 0.010	0.367 $\pm$ 0.007	0.378 $\pm$ 0.006	0.379 $\pm$ 0.008
(quad2, quad4)	0.130 $\pm$ 0.011	0.122 $\pm$ 0.009	0.119 $\pm$ 0.010	0.299 $\pm$ 0.008	0.282 $\pm$ 0.025	0.279 $\pm$ 0.008	0.403 $\pm$ 0.010	0.404 $\pm$ 0.013	0.382 $\pm$ 0.014
(quad2, size)	0.126 $\pm$ 0.009	0.109 $\pm$ 0.006	0.120 $\pm$ 0.005	0.191 $\pm$ 0.011	0.166 $\pm$ 0.010	0.181 $\pm$ 0.010	0.368 $\pm$ 0.004	0.365 $\pm$ 0.005	0.363 $\pm$ 0.004
(quad2, orientation)	0.135 $\pm$ 0.004	0.120 $\pm$ 0.010	0.133 $\pm$ 0.006	0.208 $\pm$ 0.011	0.195 $\pm$ 0.020	0.198 $\pm$ 0.010	0.368 $\pm$ 0.005	0.366 $\pm$ 0.005	0.351 $\pm$ 0.006
composition mean	0.127 $\pm$ 0.007	0.121 $\pm$ 0.008	0.134 $\pm$ 0.007	0.222 $\pm$ 0.009	0.212 $\pm$ 0.011	0.203 $\pm$ 0.010	0.377 $\pm$ 0.008	0.374 $\pm$ 0.009	0.368 $\pm$ 0.006

Table 5: Evaluation of the baselines in terms of reconstruction, concept learning, and composition on 3DIdent.

	Ada-GVAE	BetaTCVAE	TVAE
reconstruction (ELBO BPD)			
in-distribution	3.267 $\pm$ 0.001	3.261 $\pm$ 0.000	3.282 $\pm$ 0.005
out-of-distribution	3.274 $\pm$ 0.001	3.264 $\pm$ 0.001	3.291 $\pm$ 0.006
concept learning (sliced Wasserstein)			
obs	0.267 $\pm$ 0.011	0.257 $\pm$ 0.012	0.245 $\pm$ 0.003
bg	0.232 $\pm$ 0.012	0.254 $\pm$ 0.009	0.262 $\pm$ 0.006
obj	0.243 $\pm$ 0.007	0.263 $\pm$ 0.010	0.259 $\pm$ 0.007
sl	0.253 $\pm$ 0.013	0.254 $\pm$ 0.007	0.250 $\pm$ 0.009
concept learning mean	0.249 $\pm$ 0.011	0.257 $\pm$ 0.010	0.254 $\pm$ 0.006
composition (sliced Wasserstein)			
(bg, obj)	0.252 $\pm$ 0.014	0.261 $\pm$ 0.011	0.256 $\pm$ 0.010
(bg, sl)	0.227 $\pm$ 0.009	0.255 $\pm$ 0.012	0.241 $\pm$ 0.013
(obj, sl)	0.274 $\pm$ 0.010	0.247 $\pm$ 0.005	0.269 $\pm$ 0.009
composition mean	0.251 $\pm$ 0.011	0.254 $\pm$ 0.009	0.256 $\pm$ 0.011

E.1.2. EXAMPLES OF CONCEPT COMPOSITION (QUAD)

Figure 5 depicts samples generated from the context module appended to the lightweight-VAE described in Appendix D.3.

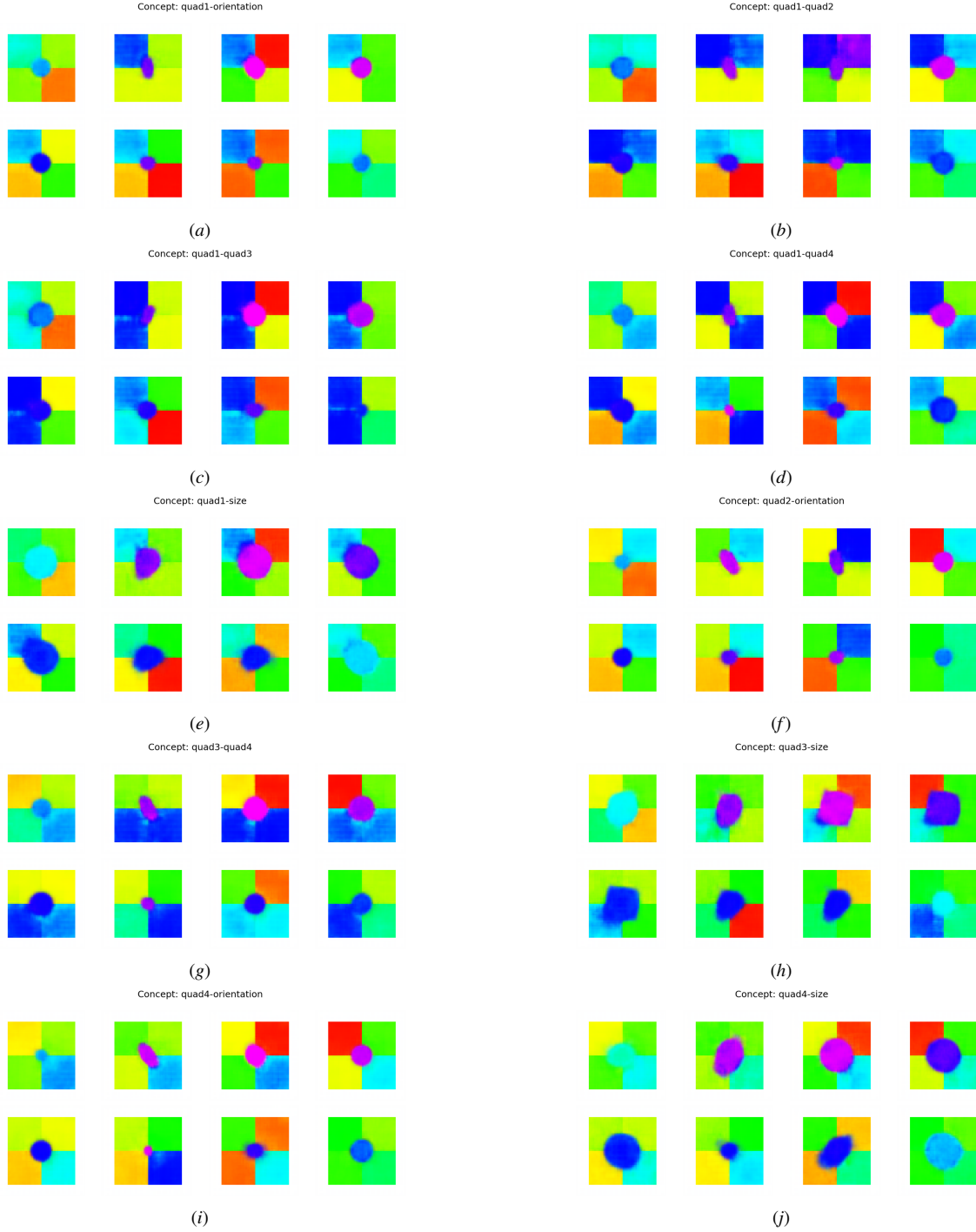


Figure 5: Example generated OOD images from a run of `quad`.

E.1.3. REGULARIZATION (MNIST)

A description of the architecture and experiment is given in Appendix D. Tables 6 and 7 respectively show results of group lasso (Yuan and Lin, 2006) and L_2 regularization, while Table 8 shows the effects regularizing the latent space of the VAE by varying β (Higgins et al., 2017).¹ Altogether, the results are relatively stable performance across regularization strengths, with slight increase in reconstruction performance for higher group lasso or L_2 regularization strength but no other statistically significant differences.

Table 6: Evaluation of CM (end-to-end) in terms of reconstruction and concept learning on MNIST for different group lasso regularization strengths.

	$\lambda = 0$	$\lambda = 1$	$\lambda = 10$	$\lambda = 100$
reconstruction (ELBO BPD)				
in-distribution	0.259 ± 0.005	0.259 ± 0.003	0.253 ± 0.003	0.248 ± 0.002
out-of-distribution	—	—	—	—
concept learning (sliced Wasserstein)				
obs	0.099 ± 0.006	0.092 ± 0.005	0.091 ± 0.006	0.092 ± 0.006
scaled	0.062 ± 0.002	0.069 ± 0.002	0.075 ± 0.003	0.118 ± 0.006
shear	0.119 ± 0.006	0.118 ± 0.005	0.115 ± 0.008	0.113 ± 0.004
shift	0.113 ± 0.005	0.112 ± 0.007	0.116 ± 0.005	0.113 ± 0.003
swel	0.126 ± 0.007	0.122 ± 0.004	0.131 ± 0.007	0.127 ± 0.006
thic	0.178 ± 0.007	0.183 ± 0.008	0.216 ± 0.015	0.220 ± 0.005
thin	0.064 ± 0.004	0.072 ± 0.005	0.082 ± 0.009	0.111 ± 0.008
concept learning mean	0.109 ± 0.006	0.110 ± 0.005	0.118 ± 0.008	0.128 ± 0.005

1. All other MNIST experiments in Appendix E.1 are run with $\lambda = 0$ for group lasso and L_2 and with $\beta = 1$.

Table 7: Evaluation of CM (end-to-end) in terms of reconstruction and concept learning on MNIST for different L_2 -norm regularization strengths.

	$\lambda = 0$	$\lambda = 1$	$\lambda = 10$	$\lambda = 100$
reconstruction (ELBO BPD)				
in-distribution	0.259 ± 0.005	0.254 ± 0.003	0.255 ± 0.004	0.253 ± 0.001
out-of-distribution	—	—	—	—
concept learning (sliced Wasserstein)				
obs	0.099 ± 0.006	0.091 ± 0.006	0.102 ± 0.005	0.086 ± 0.004
scaled	0.062 ± 0.002	0.068 ± 0.005	0.061 ± 0.003	0.073 ± 0.004
shear	0.119 ± 0.006	0.119 ± 0.007	0.127 ± 0.005	0.111 ± 0.006
shift	0.113 ± 0.005	0.111 ± 0.006	0.121 ± 0.005	0.108 ± 0.003
swel	0.126 ± 0.007	0.121 ± 0.009	0.133 ± 0.005	0.115 ± 0.004
thic	0.178 ± 0.007	0.186 ± 0.015	0.202 ± 0.010	0.183 ± 0.005
thin	0.064 ± 0.004	0.072 ± 0.008	0.062 ± 0.003	0.074 ± 0.006
concept learning mean	0.109 ± 0.006	0.110 ± 0.008	0.115 ± 0.005	0.107 ± 0.005

Table 8: Evaluation of CM (end-to-end) in terms of reconstruction and concept learning on MNIST for different latent space regularization strengths, as β varies.

	0.5	1.0	2.0	4.0
reconstruction (ELBO BPD)				
in-distribution	0.222 ± 0.002	0.259 ± 0.005	0.298 ± 0.003	0.341 ± 0.007
out-of-distribution	—	—	—	—
concept learning (sliced Wasserstein)				
obs	0.085 ± 0.003	0.099 ± 0.006	0.102 ± 0.003	0.108 ± 0.005
scaled	0.077 ± 0.006	0.062 ± 0.002	0.077 ± 0.004	0.098 ± 0.009
shear	0.108 ± 0.003	0.119 ± 0.006	0.125 ± 0.005	0.117 ± 0.004
shift	0.105 ± 0.003	0.113 ± 0.005	0.125 ± 0.004	0.130 ± 0.007
swel	0.110 ± 0.007	0.126 ± 0.007	0.129 ± 0.006	0.134 ± 0.004
thic	0.182 ± 0.011	0.178 ± 0.007	0.190 ± 0.016	0.192 ± 0.008
thin	0.069 ± 0.006	0.064 ± 0.004	0.074 ± 0.005	0.117 ± 0.005
concept learning mean	0.105 ± 0.006	0.109 ± 0.006	0.118 ± 0.006	0.128 ± 0.006

E.1.4. REGULARIZATION (QUAD)

This is like Appendix E.1.3 but applied to the quad dataset rather than MNIST. A description of the architecture and experiment is given in Appendix D. Tables 9 and 10 respectively show results of group lasso (Yuan and Lin, 2006) and L_2 regularization, while Table 11 shows the effects regularizing the latent space of the VAE by varying β (Higgins et al., 2017).²

Table 9: Evaluation of CM (end-to-end) in terms of reconstruction, concept learning, and composition on quad (independent) for different group norm regularization strengths.

	$\lambda = 0$	$\lambda = 1$	$\lambda = 10$	$\lambda = 100$
reconstruction (ELBO BPD)				
in-distribution	0.519 ± 0.022	0.511 ± 0.023	0.528 ± 0.030	0.492 ± 0.009
out-of-distribution	0.720 ± 0.069	0.703 ± 0.052	0.868 ± 0.147	0.993 ± 0.181
concept learning (sliced Wasserstein)				
obs	0.056 ± 0.003	0.049 ± 0.001	0.054 ± 0.004	0.055 ± 0.002
quad1	0.062 ± 0.003	0.058 ± 0.004	0.056 ± 0.005	0.060 ± 0.001
quad2	0.061 ± 0.005	0.057 ± 0.006	0.057 ± 0.003	0.055 ± 0.004
quad3	0.063 ± 0.002	0.058 ± 0.004	0.058 ± 0.003	0.061 ± 0.006
quad4	0.065 ± 0.007	0.054 ± 0.003	0.052 ± 0.003	0.056 ± 0.003
size	0.075 ± 0.007	0.076 ± 0.007	0.073 ± 0.007	0.069 ± 0.005
orientation	0.057 ± 0.003	0.057 ± 0.003	0.055 ± 0.006	0.065 ± 0.006
concept learning mean	0.063 ± 0.004	0.058 ± 0.004	0.058 ± 0.004	0.060 ± 0.004
composition (sliced Wasserstein)				
(quad1, quad2)	0.100 ± 0.028	0.061 ± 0.008	0.092 ± 0.023	0.101 ± 0.038
(quad1, quad3)	0.078 ± 0.016	0.080 ± 0.019	0.101 ± 0.017	0.134 ± 0.019
(quad1, quad4)	0.094 ± 0.032	0.081 ± 0.017	0.127 ± 0.040	0.132 ± 0.043
(quad1, size)	0.101 ± 0.010	0.098 ± 0.006	0.115 ± 0.011	0.127 ± 0.025
(quad1, orientation)	0.060 ± 0.003	0.055 ± 0.004	0.057 ± 0.004	0.063 ± 0.005
(quad2, quad3)	0.067 ± 0.007	0.072 ± 0.008	0.075 ± 0.008	0.123 ± 0.042
(quad2, quad4)	0.068 ± 0.004	0.087 ± 0.029	0.075 ± 0.006	0.145 ± 0.024
(quad2, size)	0.099 ± 0.008	0.097 ± 0.004	0.111 ± 0.019	0.108 ± 0.022
(quad2, orientation)	0.063 ± 0.004	0.059 ± 0.004	0.058 ± 0.004	0.065 ± 0.008
composition mean	0.081 ± 0.013	0.077 ± 0.011	0.090 ± 0.015	0.111 ± 0.025

2. All other quad experiments in Appendix E.1 are run with $\lambda = 0$ for group lasso and L_2 and with $\beta = 1$.

Table 10: Evaluation of CM (end-to-end) in terms of reconstruction, concept learning, and composition on quad (independent) for different L_2 -norm regularization strengths.

	$\lambda = 0$	$\lambda = 1$	$\lambda = 10$	$\lambda = 100$
reconstruction (ELBO BPD)				
in-distribution	0.519 ± 0.022	0.529 ± 0.029	0.550 ± 0.029	0.507 ± 0.024
out-of-distribution	0.720 ± 0.069	0.662 ± 0.058	0.676 ± 0.023	0.685 ± 0.040
concept learning (sliced Wasserstein)				
obs	0.056 ± 0.003	0.054 ± 0.001	0.055 ± 0.002	0.054 ± 0.001
quad1	0.062 ± 0.003	0.058 ± 0.003	0.061 ± 0.003	0.053 ± 0.004
quad2	0.061 ± 0.005	0.054 ± 0.003	0.062 ± 0.002	0.056 ± 0.002
quad3	0.063 ± 0.002	0.057 ± 0.002	0.059 ± 0.005	0.055 ± 0.005
quad4	0.065 ± 0.007	0.053 ± 0.003	0.058 ± 0.002	0.052 ± 0.004
size	0.075 ± 0.007	0.073 ± 0.009	0.076 ± 0.007	0.070 ± 0.004
orientation	0.057 ± 0.003	0.058 ± 0.004	0.063 ± 0.004	0.056 ± 0.005
concept learning mean	0.063 ± 0.004	0.058 ± 0.003	0.062 ± 0.003	0.057 ± 0.004
composition (sliced Wasserstein)				
(quad1, quad2)	0.100 ± 0.028	0.070 ± 0.007	0.069 ± 0.003	0.058 ± 0.003
(quad1, quad3)	0.078 ± 0.016	0.073 ± 0.014	0.096 ± 0.021	0.063 ± 0.005
(quad1, quad4)	0.094 ± 0.032	0.101 ± 0.029	0.081 ± 0.012	0.079 ± 0.019
(quad1, size)	0.101 ± 0.010	0.095 ± 0.011	0.098 ± 0.004	0.093 ± 0.007
(quad1, orientation)	0.060 ± 0.003	0.056 ± 0.004	0.061 ± 0.002	0.058 ± 0.003
(quad2, quad3)	0.067 ± 0.007	0.071 ± 0.009	0.080 ± 0.006	0.071 ± 0.009
(quad2, quad4)	0.068 ± 0.004	0.072 ± 0.009	0.067 ± 0.003	0.063 ± 0.006
(quad2, size)	0.099 ± 0.008	0.084 ± 0.006	0.097 ± 0.008	0.089 ± 0.005
(quad2, orientation)	0.063 ± 0.004	0.058 ± 0.003	0.061 ± 0.003	0.061 ± 0.003
composition mean	0.081 ± 0.013	0.076 ± 0.010	0.079 ± 0.007	0.071 ± 0.007

Table 11: Evaluation of CM (end-to-end) in terms of reconstruction, concept learning, and composition on quad (independent) for different latent space regularization strengths, as β varies.

	0.5	1.0	2.0	4.0
reconstruction (ELBO BPD)				
in-distribution	0.474 $\pm$ 0.003	0.519 $\pm$ 0.022	0.599 $\pm$ 0.020	0.640 $\pm$ 0.029
out-of-distribution	0.632 $\pm$ 0.028	0.720 $\pm$ 0.069	0.821 $\pm$ 0.047	0.865 $\pm$ 0.061
concept learning (sliced Wasserstein)				
obs	0.050 $\pm$ 0.001	0.056 $\pm$ 0.003	0.048 $\pm$ 0.002	0.054 $\pm$ 0.004
quad1	0.051 $\pm$ 0.003	0.062 $\pm$ 0.003	0.055 $\pm$ 0.002	0.064 $\pm$ 0.005
quad2	0.052 $\pm$ 0.002	0.061 $\pm$ 0.005	0.052 $\pm$ 0.003	0.064 $\pm$ 0.011
quad3	0.059 $\pm$ 0.002	0.063 $\pm$ 0.002	0.047 $\pm$ 0.002	0.064 $\pm$ 0.009
quad4	0.054 $\pm$ 0.002	0.065 $\pm$ 0.007	0.050 $\pm$ 0.002	0.058 $\pm$ 0.006
size	0.061 $\pm$ 0.003	0.075 $\pm$ 0.007	0.066 $\pm$ 0.005	0.082 $\pm$ 0.005
orientation	0.050 $\pm$ 0.002	0.057 $\pm$ 0.003	0.051 $\pm$ 0.002	0.068 $\pm$ 0.008
concept learning mean	0.054 $\pm$ 0.002	0.063 $\pm$ 0.004	0.053 $\pm$ 0.003	0.065 $\pm$ 0.007
composition (sliced Wasserstein)				
(quad1, quad2)	0.059 $\pm$ 0.005	0.100 $\pm$ 0.028	0.080 $\pm$ 0.005	0.082 $\pm$ 0.015
(quad1, quad3)	0.070 $\pm$ 0.008	0.078 $\pm$ 0.016	0.071 $\pm$ 0.007	0.082 $\pm$ 0.014
(quad1, quad4)	0.076 $\pm$ 0.010	0.094 $\pm$ 0.032	0.073 $\pm$ 0.012	0.123 $\pm$ 0.027
(quad1, size)	0.099 $\pm$ 0.007	0.101 $\pm$ 0.010	0.118 $\pm$ 0.008	0.115 $\pm$ 0.010
(quad1, orientation)	0.052 $\pm$ 0.002	0.060 $\pm$ 0.003	0.057 $\pm$ 0.003	0.073 $\pm$ 0.005
(quad2, quad3)	0.056 $\pm$ 0.000	0.067 $\pm$ 0.007	0.074 $\pm$ 0.008	0.094 $\pm$ 0.021
(quad2, quad4)	0.087 $\pm$ 0.018	0.068 $\pm$ 0.004	0.101 $\pm$ 0.027	0.067 $\pm$ 0.009
(quad2, size)	0.086 $\pm$ 0.015	0.099 $\pm$ 0.008	0.121 $\pm$ 0.018	0.131 $\pm$ 0.011
(quad2, orientation)	0.055 $\pm$ 0.003	0.063 $\pm$ 0.004	0.056 $\pm$ 0.006	0.068 $\pm$ 0.009
composition mean	0.071 $\pm$ 0.008	0.081 $\pm$ 0.013	0.083 $\pm$ 0.010	0.093 $\pm$ 0.013

E.1.5. EXPRESSIVITY (MNIST)

A description of the architecture and experiment is given in Appendix D. Tables 12 and 13 show varying widths for an expressive layer depth respectively of 2 and 5: the tuples in the column headers indicate (w_{exp}, w_c) .³ These results suggest relatively stable results across different configurations for concept learning while the best reconstruction performance is obtained from a moderate bottleneck, with expressive layer input width $w_{\text{exp}} = 22$ and a concept width of $w_c = 5$.

Table 12: Evaluation of CM (end-to-end) in terms of reconstruction and concept learning on MNIST for different capacities, with depth $h_{\text{exp}} = 2$ as (w_{exp}, w_c) varies.

(w_{exp}, w_c)	(15, 22)	(15, 5)	(22, 22)	(22, 5)	(50, 22)	(50, 5)
reconstruction (ELBO BPD)						
in-distribution	0.260 ± 0.005	0.259 ± 0.005	0.255 ± 0.003	0.252 ± 0.003	0.265 ± 0.001	0.265 ± 0.006
out-of-distribution	—	—	—	—	—	—
concept learning (sliced Wasserstein)						
obs	0.087 ± 0.003	0.099 ± 0.006	0.082 ± 0.005	0.095 ± 0.006	0.086 ± 0.004	0.092 ± 0.008
scaled	0.073 ± 0.011	0.062 ± 0.002	0.065 ± 0.004	0.067 ± 0.005	0.068 ± 0.003	0.072 ± 0.004
shear	0.115 ± 0.006	0.119 ± 0.006	0.108 ± 0.006	0.122 ± 0.006	0.116 ± 0.004	0.118 ± 0.004
shift	0.106 ± 0.009	0.113 ± 0.005	0.108 ± 0.008	0.111 ± 0.002	0.096 ± 0.003	0.115 ± 0.005
swel	0.119 ± 0.006	0.126 ± 0.007	0.107 ± 0.005	0.118 ± 0.006	0.111 ± 0.007	0.118 ± 0.007
thic	0.202 ± 0.016	0.178 ± 0.007	0.180 ± 0.007	0.171 ± 0.014	0.176 ± 0.010	0.178 ± 0.015
thin	0.065 ± 0.006	0.064 ± 0.004	0.069 ± 0.005	0.072 ± 0.010	0.073 ± 0.007	0.067 ± 0.006
concept learning mean	0.110 ± 0.008	0.109 ± 0.006	0.103 ± 0.006	0.108 ± 0.007	0.104 ± 0.005	0.109 ± 0.007

Table 13: Evaluation of CM (end-to-end) in terms of reconstruction and concept learning on MNIST for different capacities, with depth $h_{\text{exp}} = 4$ as (w_{exp}, w_c) varies.

(w_{exp}, w_c)	(15, 22)	(15, 5)	(22, 22)	(22, 5)	(50, 22)	(50, 5)
reconstruction (ELBO BPD)						
in-distribution	0.269 ± 0.003	0.272 ± 0.006	0.276 ± 0.004	0.255 ± 0.004	0.281 ± 0.005	0.279 ± 0.008
out-of-distribution	—	—	—	—	—	—
concept learning (sliced Wasserstein)						
obs	0.096 ± 0.004	0.088 ± 0.004	0.089 ± 0.001	0.092 ± 0.004	0.090 ± 0.005	0.093 ± 0.003
scaled	0.063 ± 0.003	0.071 ± 0.004	0.071 ± 0.003	0.085 ± 0.018	0.067 ± 0.001	0.079 ± 0.003
shear	0.119 ± 0.006	0.110 ± 0.004	0.116 ± 0.004	0.113 ± 0.004	0.107 ± 0.004	0.103 ± 0.004
shift	0.114 ± 0.004	0.109 ± 0.006	0.108 ± 0.005	0.107 ± 0.005	0.113 ± 0.006	0.104 ± 0.002
swel	0.120 ± 0.006	0.113 ± 0.007	0.110 ± 0.001	0.122 ± 0.006	0.116 ± 0.005	0.115 ± 0.003
thic	0.179 ± 0.012	0.178 ± 0.014	0.175 ± 0.011	0.185 ± 0.011	0.188 ± 0.007	0.184 ± 0.011
thin	0.066 ± 0.005	0.076 ± 0.004	0.072 ± 0.003	0.070 ± 0.003	0.076 ± 0.007	0.083 ± 0.006
concept learning mean	0.108 ± 0.005	0.106 ± 0.006	0.106 ± 0.004	0.111 ± 0.007	0.108 ± 0.005	0.109 ± 0.005

3. All other experiments on MNIST in Appendix E.1 are run with $h_{\text{exp}} = 2$, $w_{\text{exp}} = 15$, and $w_c = 5$.

E.1.6. EXPRESSIVITY (QUAD)

This is like Appendix E.1.5 but applied to the quad dataset rather than MNIST. A description of the architecture and experiment is given in Appendix D. Tables 14 and 15 show varying widths for an expressive layer depth respectively of 2 and 5: the tuples in the column headers indicate (w_{exp}, w_c) .⁴ These results suggest relatively stable results across different configurations.

Table 14: Evaluation of CM (end-to-end) in terms of reconstruction, concept learning, and composition on quad (independent) for different capacities, with depth $h_{\text{exp}} = 2$ as (w_{exp}, w_c) varies.

(w_{exp}, w_c)	(16, 16)	(16, 32)	(16, 8)	(32, 16)	(32, 32)	(32, 8)
reconstruction (ELBO BPD)						
in-distribution	0.494 $\pm$ 0.014	0.497 $\pm$ 0.012	0.494 $\pm$ 0.010	0.512 $\pm$ 0.023	0.557 $\pm$ 0.016	0.519 $\pm$ 0.022
out-of-distribution	0.676 $\pm$ 0.024	0.665 $\pm$ 0.025	0.598 $\pm$ 0.025	0.626 $\pm$ 0.030	0.752 $\pm$ 0.067	0.720 $\pm$ 0.069
concept learning (sliced Wasserstein)						
obs	0.055 $\pm$ 0.005	0.056 $\pm$ 0.002	0.056 $\pm$ 0.004	0.053 $\pm$ 0.003	0.063 $\pm$ 0.005	0.056 $\pm$ 0.003
quad1	0.057 $\pm$ 0.004	0.061 $\pm$ 0.005	0.059 $\pm$ 0.004	0.058 $\pm$ 0.004	0.067 $\pm$ 0.007	0.062 $\pm$ 0.003
quad2	0.061 $\pm$ 0.006	0.053 $\pm$ 0.003	0.060 $\pm$ 0.003	0.056 $\pm$ 0.005	0.063 $\pm$ 0.005	0.061 $\pm$ 0.005
quad3	0.058 $\pm$ 0.004	0.058 $\pm$ 0.005	0.065 $\pm$ 0.007	0.055 $\pm$ 0.003	0.066 $\pm$ 0.006	0.063 $\pm$ 0.002
quad4	0.060 $\pm$ 0.007	0.063 $\pm$ 0.004	0.063 $\pm$ 0.006	0.055 $\pm$ 0.003	0.068 $\pm$ 0.006	0.065 $\pm$ 0.007
size	0.072 $\pm$ 0.008	0.072 $\pm$ 0.007	0.076 $\pm$ 0.005	0.071 $\pm$ 0.005	0.082 $\pm$ 0.011	0.075 $\pm$ 0.007
orientation	0.061 $\pm$ 0.006	0.055 $\pm$ 0.002	0.058 $\pm$ 0.003	0.055 $\pm$ 0.003	0.062 $\pm$ 0.006	0.057 $\pm$ 0.003
concept learning mean	0.061 $\pm$ 0.006	0.060 $\pm$ 0.004	0.062 $\pm$ 0.005	0.057 $\pm$ 0.004	0.067 $\pm$ 0.007	0.063 $\pm$ 0.004
composition (sliced Wasserstein)						
(quad1, quad2)	0.091 $\pm$ 0.010	0.086 $\pm$ 0.010	0.066 $\pm$ 0.006	0.076 $\pm$ 0.014	0.131 $\pm$ 0.017	0.100 $\pm$ 0.028
(quad1, quad3)	0.064 $\pm$ 0.008	0.083 $\pm$ 0.006	0.075 $\pm$ 0.012	0.059 $\pm$ 0.005	0.102 $\pm$ 0.010	0.078 $\pm$ 0.016
(quad1, quad4)	0.082 $\pm$ 0.019	0.069 $\pm$ 0.006	0.064 $\pm$ 0.007	0.076 $\pm$ 0.012	0.092 $\pm$ 0.013	0.094 $\pm$ 0.032
(quad1, size)	0.084 $\pm$ 0.007	0.098 $\pm$ 0.009	0.094 $\pm$ 0.005	0.107 $\pm$ 0.013	0.117 $\pm$ 0.011	0.101 $\pm$ 0.010
(quad1, orientation)	0.059 $\pm$ 0.004	0.065 $\pm$ 0.006	0.060 $\pm$ 0.004	0.055 $\pm$ 0.003	0.068 $\pm$ 0.011	0.060 $\pm$ 0.003
(quad2, quad3)	0.083 $\pm$ 0.017	0.080 $\pm$ 0.014	0.076 $\pm$ 0.015	0.066 $\pm$ 0.008	0.083 $\pm$ 0.012	0.067 $\pm$ 0.007
(quad2, quad4)	0.079 $\pm$ 0.010	0.074 $\pm$ 0.007	0.071 $\pm$ 0.010	0.065 $\pm$ 0.009	0.077 $\pm$ 0.008	0.068 $\pm$ 0.004
(quad2, size)	0.100 $\pm$ 0.014	0.085 $\pm$ 0.011	0.087 $\pm$ 0.006	0.093 $\pm$ 0.006	0.116 $\pm$ 0.011	0.099 $\pm$ 0.008
(quad2, orientation)	0.062 $\pm$ 0.006	0.057 $\pm$ 0.003	0.060 $\pm$ 0.005	0.057 $\pm$ 0.003	0.067 $\pm$ 0.008	0.063 $\pm$ 0.004
composition mean	0.078 $\pm$ 0.011	0.077 $\pm$ 0.008	0.072 $\pm$ 0.008	0.073 $\pm$ 0.008	0.095 $\pm$ 0.011	0.081 $\pm$ 0.013

4. All other experiments on quad in Appendix E.1 are run with $h_{\text{exp}} = 2$, $w_{\text{exp}} = 32$, and $w_c = 8$.

Table 15: Evaluation of CM (end-to-end) in terms of reconstruction, concept learning, and composition on quad (independent) for different capacities, with depth $h_{exp} = 4$ as (w_{exp}, w_c) varies.

(w_{exp}, w_c)	(16, 16)	(16, 32)	(16, 8)	(32, 16)	(32, 32)	(32, 8)
reconstruction (ELBO BPD)						
in-distribution	0.511 $\pm$ 0.008	0.516 $\pm$ 0.015	0.573 $\pm$ 0.045	0.590 $\pm$ 0.009	0.613 $\pm$ 0.020	0.580 $\pm$ 0.017
out-of-distribution	0.761 $\pm$ 0.040	0.750 $\pm$ 0.033	0.674 $\pm$ 0.049	0.706 $\pm$ 0.035	0.759 $\pm$ 0.081	0.636 $\pm$ 0.017
concept learning (sliced Wasserstein)						
obs	0.059 $\pm$ 0.004	0.052 $\pm$ 0.002	0.052 $\pm$ 0.003	0.053 $\pm$ 0.002	0.059 $\pm$ 0.006	0.065 $\pm$ 0.004
quad1	0.061 $\pm$ 0.005	0.054 $\pm$ 0.004	0.061 $\pm$ 0.005	0.063 $\pm$ 0.002	0.081 $\pm$ 0.013	0.075 $\pm$ 0.005
quad2	0.057 $\pm$ 0.004	0.057 $\pm$ 0.002	0.058 $\pm$ 0.003	0.064 $\pm$ 0.004	0.080 $\pm$ 0.014	0.071 $\pm$ 0.005
quad3	0.059 $\pm$ 0.006	0.058 $\pm$ 0.003	0.064 $\pm$ 0.003	0.066 $\pm$ 0.004	0.078 $\pm$ 0.014	0.073 $\pm$ 0.002
quad4	0.056 $\pm$ 0.004	0.053 $\pm$ 0.002	0.064 $\pm$ 0.003	0.058 $\pm$ 0.004	0.075 $\pm$ 0.012	0.071 $\pm$ 0.002
size	0.069 $\pm$ 0.003	0.081 $\pm$ 0.005	0.077 $\pm$ 0.008	0.097 $\pm$ 0.004	0.111 $\pm$ 0.016	0.087 $\pm$ 0.006
orientation	0.059 $\pm$ 0.004	0.055 $\pm$ 0.003	0.061 $\pm$ 0.005	0.060 $\pm$ 0.003	0.071 $\pm$ 0.012	0.071 $\pm$ 0.003
concept learning mean	0.060 $\pm$ 0.004	0.059 $\pm$ 0.003	0.063 $\pm$ 0.004	0.066 $\pm$ 0.003	0.079 $\pm$ 0.012	0.073 $\pm$ 0.004
composition (sliced Wasserstein)						
(quad1, quad2)	0.071 $\pm$ 0.008	0.090 $\pm$ 0.010	0.059 $\pm$ 0.004	0.089 $\pm$ 0.009	0.091 $\pm$ 0.016	0.074 $\pm$ 0.005
(quad1, quad3)	0.062 $\pm$ 0.006	0.062 $\pm$ 0.004	0.080 $\pm$ 0.023	0.075 $\pm$ 0.005	0.086 $\pm$ 0.014	0.079 $\pm$ 0.006
(quad1, quad4)	0.112 $\pm$ 0.025	0.072 $\pm$ 0.010	0.063 $\pm$ 0.009	0.077 $\pm$ 0.004	0.086 $\pm$ 0.010	0.083 $\pm$ 0.005
(quad1, size)	0.083 $\pm$ 0.003	0.119 $\pm$ 0.008	0.102 $\pm$ 0.007	0.106 $\pm$ 0.003	0.133 $\pm$ 0.008	0.112 $\pm$ 0.005
(quad1, orientation)	0.059 $\pm$ 0.003	0.057 $\pm$ 0.003	0.062 $\pm$ 0.006	0.064 $\pm$ 0.003	0.079 $\pm$ 0.011	0.074 $\pm$ 0.004
(quad2, quad3)	0.077 $\pm$ 0.011	0.072 $\pm$ 0.002	0.061 $\pm$ 0.004	0.083 $\pm$ 0.008	0.109 $\pm$ 0.028	0.076 $\pm$ 0.006
(quad2, quad4)	0.092 $\pm$ 0.016	0.083 $\pm$ 0.017	0.065 $\pm$ 0.004	0.081 $\pm$ 0.005	0.102 $\pm$ 0.021	0.075 $\pm$ 0.007
(quad2, size)	0.091 $\pm$ 0.006	0.108 $\pm$ 0.009	0.101 $\pm$ 0.016	0.128 $\pm$ 0.012	0.135 $\pm$ 0.021	0.108 $\pm$ 0.007
(quad2, orientation)	0.058 $\pm$ 0.004	0.062 $\pm$ 0.003	0.064 $\pm$ 0.002	0.067 $\pm$ 0.005	0.081 $\pm$ 0.015	0.073 $\pm$ 0.005
composition mean	0.078 $\pm$ 0.009	0.080 $\pm$ 0.007	0.073 $\pm$ 0.008	0.086 $\pm$ 0.006	0.100 $\pm$ 0.016	0.084 $\pm$ 0.005

E.2. Additional NVAE Results

E.2.1. ADDITIONAL EXAMPLES OF CONCEPT COMPOSITION

Figure 6 contains additional examples of composing concepts from the NVAE-based context module. Each row corresponds to actual generated samples from the trained model in different contexts:

1. The first row shows generated observational samples;
2. The second and third row contain generated single-node interventions;
3. The final row shows generated samples where two distinct concepts are composed together.

The final row of compositional generations is genuinely OOD, in that the training data does not contain any examples where both concepts have been intervened upon.

Additional examples of both concept learning (sampling after single-target interventions) and composition (sampling after multi-target interventions) are provided in Appendices E.2.2–E.2.3.

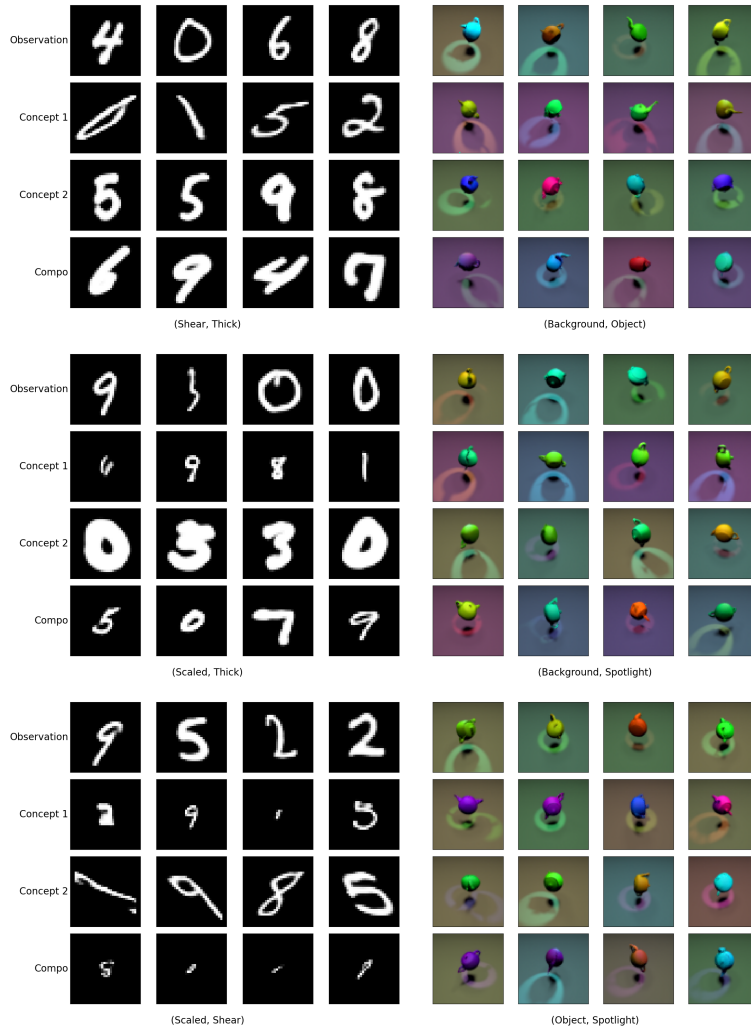


Figure 6: Additional examples of concept composition in MNIST (left) and 3DIdent (right).

E.2.2. ADDITIONAL MNIST FIGURES

See Figure 7 for concept learning and Figure 8 for composition.

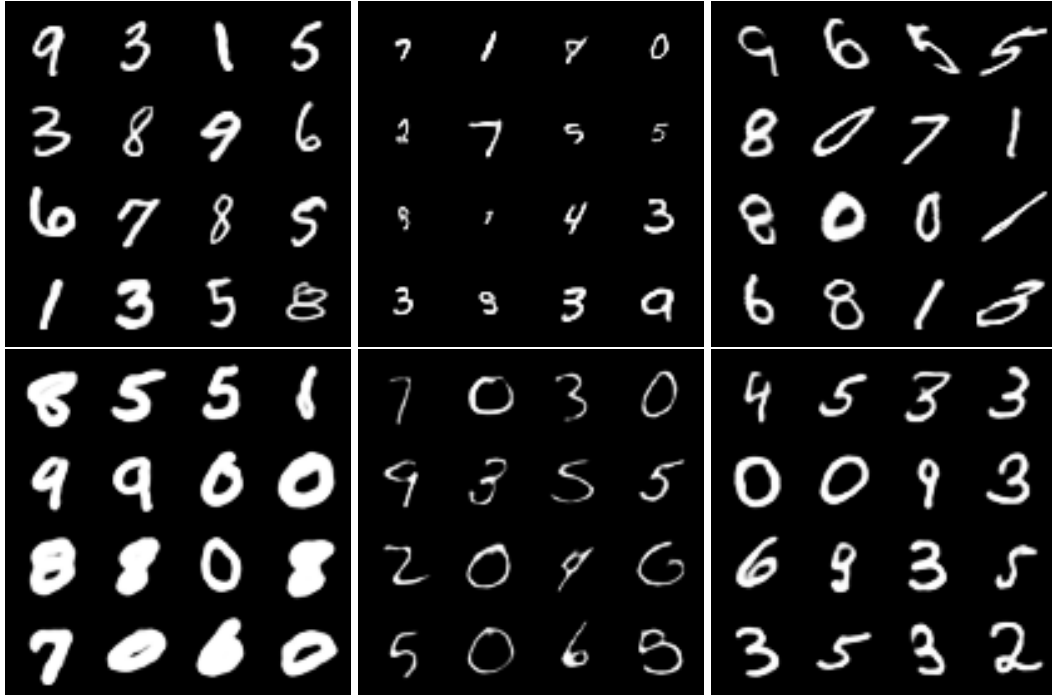


Figure 7: Examples of learned concepts (sampling after single-target intervention) in MNIST. (top) observational, scaled, shear (bottom) thic, thin, swel.



Figure 8: Examples of OOD concept composition (sampling after multi-target intervention) in MNIST.

E.2.3. ADDITIONAL 3DIDENT FIGURES

See Figure 9 for concept learning and Figure 10 for composition.



Figure 9: Examples of learned concepts (sampling after single-target intervention) in 3DIdent. (top) observational, background (bottom) object, spotlight.

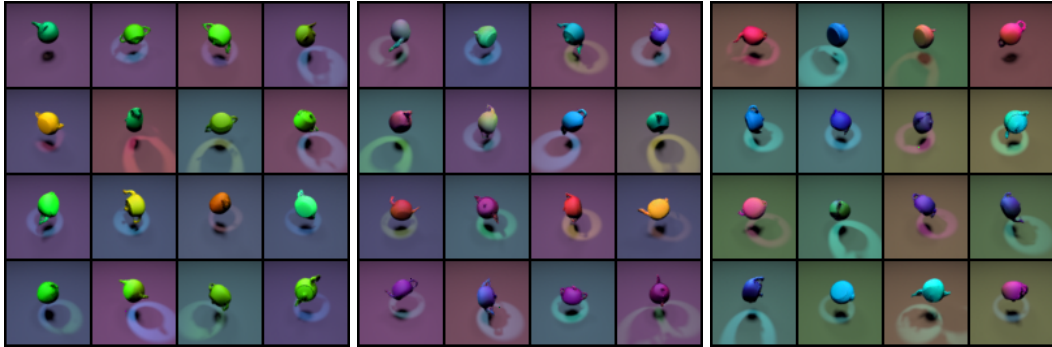


Figure 10: Examples of OOD concept composition (sampling after multi-target intervention) in 3DIdent.