
Diverse Influence Component Analysis: A Geometric Approach to Nonlinear Mixture Identifiability

Hoang-Son Nguyen and Xiao Fu

School of Electrical Engineering and Computer Science
Oregon State University
{nguyhoa3, xiao.fu}@oregonstate.edu

Abstract

Latent component identification from unknown *nonlinear* mixtures is a foundational challenge in machine learning, with applications in tasks such as self-supervised learning and causal representation learning. Prior work in *nonlinear independent component analysis* (nICA) has shown that auxiliary signals—such as weak supervision—can support *identifiability* of conditionally independent latent components. More recent approaches explore structural assumptions, e.g., sparsity in the Jacobian of the mixing function, to relax such requirements. In this work, we introduce *Diverse Influence Component Analysis* (DICA), a framework that exploits the convex geometry of the mixing function’s Jacobian. We propose a *Jacobian Volume Maximization* (J-VolMax) criterion, which enables latent component identification by encouraging diversity in their influence on the observed variables. Under reasonable conditions, this approach achieves identifiability without relying on auxiliary information, latent component independence, or Jacobian sparsity assumptions. These results extend the scope of identifiability analysis and offer a complementary perspective to existing methods.

1 Introduction

Nonlinear mixture model identification (NMMI) seeks to uncover latent components transformed by *unknown* nonlinear functions. A typical data model of interest in the context of NMMI is as follows:

$$\mathbf{x} = \mathbf{f}(\mathbf{s}), \mathbf{s} \in \mathbb{R}^d, \mathbf{x} \in \mathbb{R}^m, \quad (1)$$

where $\mathbf{s} = (s_1, \dots, s_d) \sim p(\mathbf{s})$ is a random vector following a certain distribution $p(\mathbf{s})$ with support $\mathcal{S}$, $\mathbf{f} : \mathbb{R}^d \rightarrow \mathbb{R}^m$ is an unknown, nonlinear mixing function, $\mathbf{x} = (x_1, \dots, x_m) \in \mathbb{R}^m$ is the observed data, and s_i for $i \in [d]$ and x_j for $j \in [m]$ represent the i th latent component and the j th observed feature, respectively. The function $\mathbf{f}$ is modeled a diffeomorphism that maps a latent variable to a d -dimensional Riemannian manifold $\mathcal{X}$ embedded in $\mathbb{R}^m$, where $m \geq d$ [1–3]. Given the observations (samples) of $\mathbf{x}$, NMMI amounts to recovering $\mathbf{s}$ and $\mathbf{f}$ up to certain acceptable or inconsequential ambiguities. NMMI and variants play a fundamental role in understanding many machine learning tasks, e.g., latent disentanglement [1, 4], causal representation learning [5, 6], object-centric learning [2], and self-supervised learning [7, 8].

The NMMI task is clearly ill-posed—in general, an infinite number of different $(\mathbf{f}, \mathbf{s})$ could be found from observations of $\mathbf{x}$ under (1). Hence, establishing *identifiability* of $\mathbf{f}$ and $\mathbf{s}$ becomes a central topic in NMMI. This identifiability challenge was extensively studied under the umbrella of *nonlinear independent component analysis* (nICA) [9–12]. A key take-home point is that nICA poses a much more challenging identification problem relative to ICA (where $\mathbf{f}$ is a linear system). That is, even if $s_1, \dots, s_d$ are statistically independent, the model in (1) is not identifiable [9].

In recent years, much progress has been made in understanding identifiability of (1). The line of work [1, 10–12] showed that if $s_1, \dots, s_d$ are *conditionally* independent given a certain auxiliary variable $\mathbf{u}$ (which can be understood as additional side information), then $\mathbf{f}$ and $\mathbf{s}$ can be recovered to reasonable extents. Another line of work tackles identifiability by assuming that the mixing function $\mathbf{f}$ is *structured* other than completely unknown. For example, the works [13–16] make an explicit structural assumption that $\mathbf{f}$ is a post-nonlinear mixing function, the work [17] assumes that $\mathbf{f}$ is conformal, and the more recent work [18] assumes that $\mathbf{f}$ is piecewise linear. Using these structures, the auxiliary information can be circumvented for establishing identifiability.

More recent advances propose to exploit *implicit* structures (other than *explicit* structures like post-nonlinearity) of $\mathbf{f}$. Notably, the works [3, 19, 20] utilize structures defined over the Jacobian of $\mathbf{f}$ for identifying the model (1). In particular, it was shown that as long as the Jacobian of $\mathbf{f}$ follows certain sparsity patterns, then the model (1) is identifiable [2, 3, 21]. Imposing structural constraints on the Jacobian of $\mathbf{f}$ is arguably less restrictive compared to assuming explicit parameterization of $\mathbf{f}$. Structures of Jacobian reflect how the latent variables $s_1, \dots, s_d$ affect the observed features $x_1, \dots, x_m$, making the related assumptions admit intuitive physical meaning. Using structured Jacobian to model such influences is also advocated by several causal representation learning paradigms (see, e.g., *independent mechanism analysis* (IMA) [20] under the *independent causal mechanism* (ICM) principle [22]), and other frameworks, e.g. object-centric representation learning [2, 23, 24].

Open Question. The advancements have been encouraging, yet understanding to NMMI identifiability remains to be deepened. In particular, the Jacobian sparsity-based approaches, e.g., [2, 3, 19, 21], provided intriguing ways of establishing identifiability of the model in (1)—without using auxiliary variables, latent component independence, or relatively restrictive structural constraints of $\mathbf{f}$. However, the Jacobian sparsity assumptions in these works essentially assume that the observed features are only generated from a subset of $s_1, \dots, s_d$, which sometimes may not hold. It is tempting to circumvent using such strict sparsity-based assumptions in NMMI.

Contributions. In this work, we propose to make use of an alternative condition of $\mathbf{f}$'s Jacobian for NMMI. We leverage the fact that, as long as the influences of $\mathbf{s}$ imposed on each x_j for $j = 1, \dots, m$ are *sufficiently diverse*, $\mathbf{f}$'s Jacobian exhibits an interesting convex geometry. This geometry is similar to the classical “*sufficiently scattered condition* (SSC)” in the *structured matrix factorization* (SMF) literature [25–29] and, particularly, the more recent developments in *polytopic matrix factorization* (PMF) [30] (also see [31–33]). As a consequence, taking intuition from volume-regularized SMF [25, 30–36], we show that fitting the data model in (1) together with maximizing the learned $\mathbf{f}$'s Jacobian volume *provably* recovers $\mathbf{f}$ and $\mathbf{s}$ up to acceptable ambiguities. The proposed *Jacobian volume maximization* (J-VolMax) approach does not rely on auxiliary variables or statistical independence. More notably, an $\mathbf{f}$ with a *dense* Jacobian can still be provably identified under our formulation. These advantages make our identifiability results applicable to a wide range of scenarios that are not covered by the existing literature, substantially expanding the understanding of model identifiability under (1). We tested the proposed approach on synthetic data and in a single-cell transcriptomics application. The results corroborate with our NMMI theory.

Notation. Please refer to Appendix A.1 for details.

2 Background

From ICA to nICA. ICA is arguably one of the most influential latent component identification approaches. Classic ICA [37, 38] assumes that $\mathbf{x} = \mathbf{A}\mathbf{s}$ with a nonsingular $\mathbf{A}$ and that

$$p(\mathbf{s}) = \prod_{i=1}^d p_i(s_i), \quad (2)$$

where at most one s_i is a Gaussian variable. Then, the identifiability of $(\mathbf{A}, \mathbf{s})$ up to permutation and scaling ambiguities can be established by finding a linear filter to output independent estimates. Unfortunately, the nICA model, i.e., (1) with condition (2), is in general non-identifiable [9].

nICA with Auxiliary Information. More recent breakthroughs propose to use the following conditional independence model, i.e.,

$$p(\mathbf{s}|\mathbf{u}) = \prod_{i=1}^d p_i(s_i|\mathbf{u}), \quad (3)$$

to establish identifiability of (1). The side information $\mathbf{u}$ can be time frame labels [10, 11, 21, 39, 40], observation group indices [41], or view indices [42]. The takeaway from this line of work is that, using

the *variations* of $\mathbf{u}$, one can fend against negative effects (e.g., the existence of measure-preserving automorphism (MPA) [6, 43]) leading to non-identifiability under (1). The results are elegant and inspiring. Nonetheless, both $\mathbf{u}$ and conditional independence might not always be available.

nICA with Structured $\mathbf{f}$. Another route to establish identifiability is to exploit prior structural information of $\mathbf{f}$. For example, the works [9, 16, 18, 23, 44–48] used conformal, local isometry, close-to-linear, post-nonlinear, piecewise affine structures, and additive structures of $\mathbf{f}$, respectively. Recent works, e.g., [14, 15, 18, 23, 48], also showed that some of these structures can be used for dependent component identification. Nonetheless, such *explicit* structures of $\mathbf{f}$, e.g., post-nonlinear mixtures, only make sense when they are used for suitable applications (e.g., hyperspectral imaging [15] and audio separation [16]), yet most problems in generative model learning and representation learning may not have such structures.

Exploiting Jacobian Structures. Instead of imposing explicit structures directly on $\mathbf{f}$, it is also plausible to exploit the structures of the Jacobian of $\mathbf{f}$. Note that $[\mathbf{J}_{\mathbf{f}}(\mathbf{s})]_{i,j} = \partial x_i / \partial s_j$ characterizes how x_i is influenced by the change of s_j .

The aforementioned IMA approach [20], inspired by the principle of ICM [22], assumes orthogonal columns of $\mathbf{J}_{\mathbf{f}}(\mathbf{s})$ at each point $\mathbf{s} \in \mathcal{S}$. But these developments lack comprehensive identifiability characterizations (see [44]). On the other hand, the works [2, 3, 19, 49] assume that $\mathbf{J}_{\mathbf{f}}(\mathbf{s})$ exhibits a certain type of sparsity pattern. These works formulate the NMMI problem as Jacobian sparsity-regularized data fitting problems. For example, the work [2] proposed the following:

$$\min_{\mathbf{f}, \mathbf{g}} \mathbb{E}_{\mathbf{x}} [\|\mathbf{f}(\mathbf{g}(\mathbf{x})) - \mathbf{x}\|_2^2 + \lambda_{\text{sp}} c_{\text{sp}}(\mathbf{J}_{\mathbf{f}}(\mathbf{g}(\mathbf{x}))), \quad (4)$$

where the first term finds a diffeomorphism $\mathbf{f}$ and its “inverse” $\mathbf{g}$ (in which $\mathbf{g}(\mathbf{x})$ is supposed to recover $\mathbf{s}$), and the second term $c_{\text{sp}}(\cdot)$ promotes sparsity of $\mathbf{J}_{\mathbf{f}}$ —see [2, 3, 19, 49, 50] for their respective ways of sparsity promotion. The most notable feature of this line of work lies in its relatively relaxed conditions on $\mathbf{s}$ for establishing identifiability of the model. For instance, the work [2] showed that identifiability can be established without using auxiliary variables or statistical independence of $\mathbf{s}$. On the other hand, $\mathbf{J}_{\mathbf{f}}$ being sparse means that the observed features in $\mathbf{x}$ are only generated by subsets of $\mathbf{s}$. This assumption is justifiable in some applications, e.g., object-centric image generation [2, 23, 24], but can be violated in other settings.

3 Proposed Approach: Diverse Influence Component Analysis

We present an alternative way to establish identifiability of the nonlinear mixture model in (1), without relying on statistical assumptions of $\mathbf{s}$, sparsity of $\mathbf{J}_{\mathbf{f}}$, or auxiliary information. Our approach is built upon *convex geometry* of $\mathbf{J}_{\mathbf{f}}(\mathbf{s})$ and an underlying connection between NMMI and SMF models [25–29], particularly the recent advancements in [30] and variants in [31–33].

Preliminaries of Convex Geometry. In Appendix A.2, we give a brief introduction to the notions (e.g., convex hull, polar set, *maximal volume inscribed ellipsoid* (MVIE)) that we use in our context.

3.1 Sufficiently Diverse Influence

Motivation. We are interested in understanding how the influences of $\mathbf{s}$ on $\mathbf{x}$ affect the identifiability of (1), beginning with a close examination of existing Jacobian assumptions. First, the sparsity assumptions on $\mathbf{J}_{\mathbf{f}}$ in [2, 3, 19, 49–51] embody key principles in generative modeling and causal representation learning. For instance, the ICM principle [22] promotes generative models where latent variables exert distinct influences on observed features [20]. Sparse Jacobian models share this view when the sparsity patterns of $[\mathbf{J}_{\mathbf{f}}]_{i,:}$ for $i = 1, \dots, m$ (or $[\mathbf{J}_{\mathbf{f}}]_{j,:}$ for $j = 1, \dots, d$) differ sufficiently. However, sparsity is arguably a relatively stringent assumption—the hard constraint that each x_i depends only on a subset of $s_1, \dots, s_d$ may not always hold in practice. Second, the IMA method [52] reflects the ICM principle differently: by enforcing column orthogonality in $\mathbf{J}_{\mathbf{f}}$, it assumes that the influences of s_i and s_j on the change of $\mathbf{x}$ are mutually perpendicular. This circumvents sparsity-based constraints and permits all x_i to depend on all latent components. Nonetheless, IMA only guarantees local identifiability of (1) (see [44]), while global identifiability remains unresolved. These observations motivate us to develop an alternative framework that flexibly models diverse interactions between $\mathbf{s}$ and $\mathbf{x}$, while ensuring global identifiability of nonlinear mixtures.

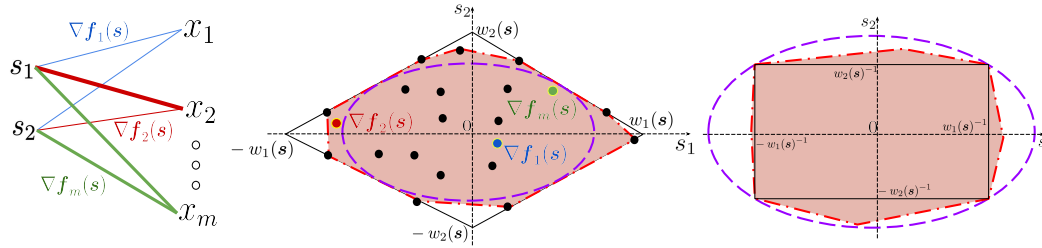


Figure 1: **[Left]** $\mathbf{s}$, $\mathbf{x}$, and $\nabla f_i(\mathbf{s}) \in \mathbb{R}^d$ for $d = 2, \forall i \in [m]$; line thickness indicates the magnitude of influence of s_i on x_j . **[Middle]** Condition 1 in Assumption 3.1 for $d = 2$: axes represent $\partial f_i / \partial s_1$ and $\partial f_i / \partial s_2 \in \mathbb{R}$; the pink region is $\text{conv}\{\nabla f_1(\mathbf{s}), \dots, \nabla f_m(\mathbf{s})\}$, the dashed purple ellipse is $\mathcal{E}(\mathcal{B}_1^{w(\mathbf{s})})$, and the solid black diamond is $\mathcal{B}_1^{w(\mathbf{s})}$. **[Right]** Condition 2 in Assumption 3.1: shaded region shows $\text{conv}\{\nabla f_1(\mathbf{s}), \dots, \nabla f_m(\mathbf{s})\}^*$, dashed ellipse shows $\mathcal{E}(\mathcal{B}^{w(\mathbf{s})})$, and solid rectangle shows $(\mathcal{B}_1^{w(\mathbf{s})})^* = \mathcal{B}_{\infty}^{w(\mathbf{s})}$.

Diverse Influence. To formally underpin the intuitive notion of “diverse influence” to an exact definition, we use convex geometry to characterize $\nabla f_1(\mathbf{s}), \dots, \nabla f_m(\mathbf{s})$ (i.e., the rows of $\mathbf{J}_f$). Note that when the rows of $\mathbf{J}_f$ are sufficiently distinct, it means that $\mathbf{s}$ ’s changes render different variations of $x_1, \dots, x_m$. In particular, we will exploit the cases where some observed dimensions $x_i = f_i(\mathbf{s})$ for $i \in [m]$ are affected by the changes of $s_1, \dots, s_d$ in a sufficiently different way.

Assumption 3.1 (Sufficiently Diverse Influence (SDI)). At $\mathbf{s} \in \mathcal{S}$, there exists an $\mathbf{s}$ -dependent weighted L_1 norm ball $\mathcal{B}_1^{w(\mathbf{s})}$ such that $\nabla f_1(\mathbf{s}), \dots, \nabla f_m(\mathbf{s}) \in \mathcal{B}_1^{w(\mathbf{s})}$. In addition, the following two conditions hold:

1. $\mathcal{E}(\mathcal{B}_1^{w(\mathbf{s})}) \subseteq \text{conv}\{\nabla f_1(\mathbf{s}), \dots, \nabla f_m(\mathbf{s})\} \subseteq \mathcal{B}_1^{w(\mathbf{s})}$, and
2. $\text{conv}\{\nabla f_1(\mathbf{s}), \dots, \nabla f_m(\mathbf{s})\}^* \cap \text{bd}(\mathcal{E}(\mathcal{B}_1^{w(\mathbf{s})})^*) = \text{extr}(\mathcal{B}_{\infty}^{w(\mathbf{s})})$,

where $\text{conv}\{\cdot\}$ denotes the convex hull of a set of vectors, and for a given polytope $\mathcal{P}$, $\mathcal{E}(\mathcal{P})$ is its MVIE, $\mathcal{P}^*$ stands for its polar set, $\text{extr}(\mathcal{P})$ contains its extreme points, and $\text{bd}(\mathcal{P})$ is the polytope’s boundary.

The SDI condition describes the scattering geometry of $\nabla f_1(\mathbf{s}), \dots, \nabla f_m(\mathbf{s})$, which originates from the *sufficiently scattered condition* (SSC) in the matrix factorization literature. The SSC first appeared in identifiability analysis of *nonnegative matrix factorization* (NMF) [26] and has various forms in the SMF literature; see [25, 27–30, 36]. Here, SDI takes the form of SSC from *polytopic matrix factorization* (PMF) [30]; also see [31–33]. The key difference between SDI and SSC is that the former specifies the row-scattering pattern of a Jacobian $\mathbf{J}_f(\mathbf{s})$ at every $\mathbf{s}$ over a continuous manifold $\mathcal{S}$, yet SSC only characterizes the latent factors of a data matrix (e.g., $\mathbf{W}, \mathbf{H}$ in $\mathbf{X} = \mathbf{W}\mathbf{H}$) that do not involve nonlinear functions or derivatives. But their geometries are essentially the same.

Geometry of SDI. Fig. 1 [Middle] illustrates Condition 1 of Assumption 3.1 for $d = 2$. One can see that SDI is about the spread of $\{\nabla f_1(\mathbf{s}), \dots, \nabla f_m(\mathbf{s})\}$ in an $\mathbf{s}$ -dependent weighted L_1 -norm ball $\mathcal{B}_1^{w(\mathbf{s})}$. We note that $\mathcal{B}_1^{w(\mathbf{s})}$ always exists, as long as $\{\nabla f_1(\mathbf{s}), \dots, \nabla f_m(\mathbf{s})\}$ are finite. Geometrically, Condition 1 requires that the gradient vectors’ convex hull contains the MVIE of $\mathcal{B}_1^{w(\mathbf{s})}$. Condition 2 is imposed on the polar sets $\text{conv}\{\nabla f_1(\mathbf{s}), \dots, \nabla f_m(\mathbf{s})\}^*$, $\mathcal{E}(\mathcal{B}_1^{w(\mathbf{s})})^*$ and $\mathcal{B}_{\infty}^{1/w(\mathbf{s})}$ (note that $\mathcal{B}_{\infty}^{1/w(\mathbf{s})}$, weighted by $1/w_1(\mathbf{s}), \dots, 1/w_d(\mathbf{s})$, is the polar of $\mathcal{B}_1^{w(\mathbf{s})}$). The condition implies that $\text{conv}\{\nabla f_1(\mathbf{s}), \dots, \nabla f_m(\mathbf{s})\}^*$ touches the boundary of $\mathcal{E}(\mathcal{B}_1^{w(\mathbf{s})})^*$ at the vertices of $\mathcal{B}_{\infty}^{1/w(\mathbf{s})}$, which is a vector of the form $[\pm w_1(\mathbf{s})^{-1}, \dots, \pm w_d(\mathbf{s})^{-1}]^T \in \mathbb{R}^d$. In the original domain, this means that the ellipsoid $\mathcal{E}(\mathcal{B}_1^{w(\mathbf{s})})$ must touch the convex hull $\text{conv}\{\nabla f_1(\mathbf{s}), \dots, \nabla f_m(\mathbf{s})\}$ at the facets of $\mathcal{B}_1^{w(\mathbf{s})}$. Fig. 1 [Right] shows the polar set view of Fig. 1 [Middle]. Readers are referred to [30] for a visualization of SSC under PMF, which further clarifies the connection between SSC and SDI.

Physical Meaning of SDI. SDI reflects how $\mathbf{s}$ diversely affects $x_1, \dots, x_m$. Under SDI, there exist some observed features positively influenced by s_j (i.e., $\partial x_i / \partial s_j > 0$), while also some features negatively influenced by s_j (i.e., $\partial x_i / \partial s_j < 0$). This pattern likely holds in many applications with high-dimensional data (i.e. $m \gg d$). This is because each x_i can only either be positively or negatively influenced by s_j . As m increases there would be more features that are affected differently by s_j .

Note that the positive and negative influences need not be symmetric (i.e., $\partial x_i / \partial s_j \neq -\partial x_{i'} / \partial s_j$ is allowed), as illustrated in Fig. 1 [Middle].

In addition, SDI assumes that each latent component s_k *dominantly* influences at least two observed features, $x_{i_k^+}$ and $x_{i_k^-}$; here, “+” and “-” indicate that the features are positively and negatively affected by s_k , respectively. In other words, features satisfying the following should exist:

$$\partial x_{i_k^+} / \partial s_k \gg \partial x_i / \partial s_j, \forall j \neq i_k^+ \quad \text{and} \quad \partial x_{i_k^-} / \partial s_k \ll \partial x_i / \partial s_j, \forall j \neq i_k^-. \quad (5)$$

Discussion. Beyond the basic geometric and physical interpretations, some additional remarks on other aspects of SDI are as follows:

(i) *Dependent s and Dense Jacobian Can Satisfy SDI.* Notably, SDI does not impose statistical assumptions on s (e.g., conditionally independent given auxiliary variables like in [10–12]). In addition, SDI does not rely on the sparsity of $\mathbf{J}_f$ as in [2, 3, 19]. In fact, $\mathbf{J}_f$ can be *completely dense* and at the same time satisfies SDI (see Fig. 1 [Middle]).

(ii) *SDI Favors $m \gg d$ Cases.* The SDI condition is arguably easier to satisfy when $m \gg d$ (which is justified in many applications, such as image/video generation [2, 40, 53]). It is evident that SDI requires at least $m = 2d$, which corresponds to the case where $\nabla f_1(s), \dots, \nabla f_m(s)$ are located at $\pm w_1(s)e_1, \dots, \pm w_d(s)e_d$. To see why SDI is in favor of larger m ’s, consider a case where the rows of $\mathbf{J}_f(s)$, i.e., $\nabla f_i(s)$, are drawn from a certain continuous distribution supported on $\mathcal{B}_1^{w(s)}$. Then, a large m means that more realizations of $\nabla f_i(s)$ ’s are available. Therefore, for a fixed d , $m \rightarrow \infty$ leads to $\text{conv}\{\nabla f_1(s), \dots, \nabla f_m(s)\}$ increasingly covering more of $\mathcal{B}_1^{w(s)}$ and eventually becoming $\mathcal{B}_1^{w(s)}$, which ensures that SDI condition is met. A related note is that a matrix factor $\mathbf{W} \in \mathbb{R}^{m \times d}$ with larger m under fixed d would have higher probabilities to satisfy SSC, which was formally studied in [54], using exactly the same insight.

(iii) *SDI Encodes Influence Variations.* Lastly, the SDI condition allows the gradient pattern to vary from point to point on $\mathcal{S}$: at different $s \in \mathcal{S}$, both the ball $\mathcal{B}_1^{w(s)}$ and the position pattern of $\nabla f_1(s), \dots, \nabla f_m(s)$ in $\mathcal{B}_1^{w(s)}$ can be different, as long as the gradients are sufficiently spread out in different directions as in Assumption 3.1.

3.2 Proposed Learning Criterion

Similar to the NMMI literature, e.g., [3, 10–12, 19], the goal of identifiability-guaranteed learning in this work is to find an invertible function (over the $\mathbb{R}^d$ manifold) $\hat{\mathbf{f}}$ such that $\hat{\mathbf{s}} = \hat{\mathbf{f}}^{-1}(\mathbf{x}) = \hat{\mathbf{f}}^{-1} \circ \mathbf{f}(s)$ recovers s to a reasonable extent. In practice, we use a neural network (supposedly a universal function representer) $\mathbf{f}_\theta$ as our learning function. To ensure the invertibility of $\mathbf{f}_\theta : \mathbb{R}^d \rightarrow \mathbb{R}^m$ over the manifold $\mathcal{S} \subseteq \mathbb{R}^d$, we use another neural network $\mathbf{g}_\phi : \mathbb{R}^m \rightarrow \mathbb{R}^d$ and enforce

$$\mathbf{x} = \mathbf{f}_\theta(\mathbf{g}_\phi(\mathbf{x})), \forall \mathbf{x} \in \mathcal{X}. \quad (6)$$

Note that if the above holds, both $\mathbf{f}_\theta$ and $\mathbf{g}_\phi$ are invertible over the data-generating d -dimensional manifold. Eq. (6) is nothing but a stacked autoencoder, which by itself does not ensure identifiability of the model (1). To utilize SDI for establishing identifiability, we propose the following learning criterion, namely, *Jacobian volume maximization* (J-VolMax):

$$\text{(J-VolMax)} \quad \underset{\theta, \phi}{\text{maximize}} \quad \mathbb{E}[\log \det(\mathbf{J}_{\mathbf{f}_\theta}(\mathbf{g}_\phi(\mathbf{x}))^\top \mathbf{J}_{\mathbf{f}_\theta}(\mathbf{g}_\phi(\mathbf{x})))] \quad (7a)$$

$$\text{subject to: } \|\mathbf{J}_{\mathbf{f}_\theta}(\mathbf{g}_\phi(\mathbf{x}))_{i,:}\|_1 \leq C, \quad \forall i = 1, \dots, m, \quad (7b)$$

$$\mathbf{x} = \mathbf{f}_\theta(\mathbf{g}_\phi(\mathbf{x})), \quad \forall \mathbf{x} \in \mathcal{X} \quad (7c)$$

The term $\log \det(\mathbf{J}_{\mathbf{f}_\theta}(\mathbf{g}_\phi(\mathbf{x}))^\top \mathbf{J}_{\mathbf{f}_\theta}(\mathbf{g}_\phi(\mathbf{x})))$ represents the squared volume of the convex hull spanned by the columns of $\mathbf{J}_{\mathbf{f}_\theta}$.

Using this criterion, the partial derivatives contained in $\mathbf{J}_{\mathbf{f}_\theta}(\hat{\mathbf{s}})$ (where $\hat{\mathbf{s}} = \mathbf{g}_\phi(\mathbf{x})$) are encouraged to scatter in space—reflecting the belief that the influences of s on different x_i ’s are diverse.

Identifiability Result. Under the model in (1) and Assumption 3.1, we show our main result:

Theorem 3.2 (Identifiability of J-VolMax). *Denote any optimal solution of Problem (7) as $(\hat{\theta}, \hat{\phi})$. Assume $\hat{f} = f_{\hat{\theta}}$ and $\hat{g} = g_{\hat{\phi}}$ are universal function representers. Suppose the model in (1) and Assumption 3.1 hold for every $s \in \mathcal{S}$. Then, we have $\hat{s} = \hat{g}(x) = \hat{g} \circ f(s)$ where*

$$[\hat{s}]_i = [\hat{g}(x)]_i = \rho_i(s_{\pi(i)}), \forall i \in [d], \quad (8)$$

in which π is a permutation of $\{1, \dots, d\}$ and $\rho_i(\cdot) : \mathbb{R} \rightarrow \mathbb{R}$ is an invertible function.

Notice that we used the constraint $\|J_{f_{\theta}}(g_{\phi}(x))_{i,:}\|_1 \leq C$ in (7) where $C > 0$ is unknown. Under SDI, the ground-truth Jacobian is bounded by $\|J_f(s)_{i,:}\|_1 \in \mathcal{B}_1^{w(s)}$ with an *unknown* $\mathcal{B}_1^{w(s)}$. Nonetheless, using an arbitrary L_1 -norm ball with radius C in (7) does not affect identifiability of the J-VolMax criterion—the learned $\hat{s}$ will have a scaling ambiguity anyway.

The proof of the theorem consists of three major steps. First, we recast the nonlinear identifiability problem J-VolMax into its first-order derivative domain as a matrix-finding problem at each $s \in \mathcal{S}$, which is closely related to intermediate steps in matrix factor identification under the PMF model [30]. Second, consequently, we utilize SDI and algebraic properties from volume maximization-based PMF to underpin identifiability under J-VolMax at each s up to an s -dependent permutation ambiguity. Third, we invoke continuity of f and its domain $\mathcal{S}$ to unify permutation ambiguity over the entire $\mathcal{S}$. The details can be found in Appendix B.

Finite-Sample SDI and Identifiability. In Theorem 3.2, an assumption is that every s in the continuous domain $\mathcal{S}$ has to satisfy SDI, which could be relatively restrictive. To proceed, we show that, as f and the learned functions satisfy certain conditions, having a finite number of s satisfying SDI suffices to establish identifiability up to bounded errors.

To see how we approach this, consider a finite set $\mathcal{S}_N := \{s^{(1)}, \dots, s^{(N)}\}$ with $\mathcal{X}_N := \{x \in \mathcal{X} : x = f(s), \forall s \in \mathcal{S}_N\}$ such that Assumption 3.1 is satisfied at each of the N points in $\mathcal{S}_N$. Note that at each $s^{(n)}$, the optimal encoder $\hat{g}$ recovers $s^{(n)}$ from the observation $x^{(n)}$ up to permutation $\hat{\Pi}(s^{(n)})$ and an invertible element-wise map $\hat{\rho}(s^{(n)})$, i.e.,

$$\hat{g}(x^{(n)}) = \hat{\Pi}(s^{(n)})\hat{\rho}(s^{(n)}), \forall x^{(n)} \in \mathcal{X}_N, \quad (9)$$

which is a direct result when using the J-VolMax learning criterion; see Lemma C.2. The following result shows that under certain regularity conditions, if the set $\mathcal{S}_N$ contains samples that locate densely enough in space, the learned encoder can recover the ground-truth s up to the same ambiguities as in Theorem 3.2, with a bounded error that decays as N grows.

Theorem 3.3 (Identifiability under Finite-sample SDI). *Assume that there is a finite set $\mathcal{S}_N := \{s^{(1)}, \dots, s^{(N)}\}$ with $\mathcal{X}_N := \{x \in \mathcal{X} : x = f(s), \forall s \in \mathcal{S}_N\}$ such that the Assumption 3.1 is satisfied at each of the N points in $\mathcal{S}_N$. Let $\hat{g} \in \mathcal{G}$ be the optimal encoder by J-VolMax criterion (7), and $\bar{\Theta}$ contains the parameters of the learned encoder and decoder. Further assume that the following regularity conditions hold:*

1. *The functions $g = f^{-1}$ and g_{ϕ} are from classes $\mathcal{G}'$ and $\mathcal{G}$, respectively, where $\mathcal{G}' \subseteq \mathcal{G}$.*
2. *The functions $f, \hat{g}, \hat{\rho}$ are Lipschitz continuous with constants $L_f, L_{\hat{g}}, L_{\hat{\rho}} > 0$.*
3. *There is a $\gamma > 0$ such that for any permutation matrix $\Pi \in \mathcal{P}_d$ and $\Pi \neq \hat{\Pi}(s^{(n)})$ (from (9)),*

$$\|\hat{g}(x^{(n)}) - \Pi\hat{\rho}(s^{(n)})\|_2 \geq \gamma, \forall n \in [N].$$

4. *For $\mathcal{N}^{(n)} = \{s \in \mathcal{S} : \|s - s^{(n)}\|_2 < r^{(n)}\}$ with $r^{(n)} < \frac{\gamma}{2(L_f L_{\hat{g}} + L_{\hat{\rho}})}$, the union of the neighborhoods, $\mathcal{N} := \bigcup_{n=1}^N \mathcal{N}^{(n)}$, is a connected subset of $\mathcal{S}$ and $V(s; \bar{\Theta})$ is optimal for any $s \in \mathcal{N}$.*
5. *The points $s^{(1)}, \dots, s^{(N)} \in \mathcal{S}_N$ densely locate in $\mathcal{S}$ such that*

$$\mathbb{P}(s \in \mathcal{N}) / \mathbb{P}(s \in \mathcal{S} \setminus \mathcal{N}) > G_{\max} / G_{\min} > 1, \quad (10)$$

where $G_{\min}, G_{\max}$ are bi-Lipschitz constants of the Jacobian volume surrogate: for any parameters Θ_1, Θ_2 in $(\mathcal{F}, \mathcal{G})$,

$$G_{\min} \|\Theta_1 - \Theta_2\|_2 \leq |V(s; \Theta_1) - V(s; \Theta_2)| \leq G_{\max} \|\Theta_1 - \Theta_2\|_2, \forall s \in \mathcal{S}. \quad (11)$$

Then, $\widehat{\mathbf{g}}(\mathbf{x}^{(n)}) = \widehat{\Pi}\widehat{\rho}(\mathbf{s}^{(n)})$, $\forall n \in [N]$ for a constant permutation matrix $\widehat{\Pi} \in \mathcal{P}_d$. Furthermore, with probability at least $1 - \delta$,

$$\mathbb{E}_{\mathbf{s} \sim p(\mathbf{s})}[\|\widehat{\mathbf{g}}(\mathbf{x}) - \widehat{\Pi}\widehat{\rho}(\mathbf{s})\|_2] = \mathcal{O}\left((L_{\mathbf{f}}L_{\widehat{\mathbf{g}}} + L_{\widehat{\rho}})\mathcal{R}_N(\mathcal{G}) + \sqrt{\ln(1/\delta)/N}\right), \quad (12)$$

where $\mathcal{R}_N(\mathcal{G})$ is the empirical Rademacher complexity of the encoder class.

The theorem implies that if the number of $\{\mathbf{s}^{(n)}\}_{n=1}^N$ satisfying SDI is sufficiently large, the identification error defined in (12) can be made arbitrarily small. Note that when $\mathcal{G}$ is not a universal function class—that is, its capacity is controlled, e.g., through bounded norms or Lipschitz constants—its empirical Rademacher complexity satisfies $\mathcal{R}_N(\mathcal{G}) = \mathcal{O}(1/\sqrt{N})$ [55]. Consequently, with probability at least $1 - \delta$, the overall error rate is $\mathcal{O}(\sqrt{\ln(1/\delta)} + 1/\sqrt{N}) = \tilde{\mathcal{O}}(1/\sqrt{N})$. The detailed proof is provided in Appendix C.

In Theorem 3.3, Condition 1 means that the learning function class can faithfully represent $\mathbf{g}$. Condition 2 requires that $\mathbf{f}$ and the learned functions have bounded continuity constants, which is a mild assumption as the learning functions are often represented using continuous function classes (e.g., neural networks). Condition 3 assumes that at each point $\mathbf{s}^{(n)} \in \mathcal{S}_N$, there is a unique optimal permutation between $\mathbf{s}^{(n)}$ and $\widehat{\mathbf{g}}(\mathbf{x}^{(n)})$. That is, $\gamma = 0$ holds for the optimal permutation, but an error of $\gamma > 0$ occurs for all other permutations. Such a condition naturally holds under sufficient continuity of the functions. Condition 4 translates to the fact that the points $\mathbf{s}^{(1)}, \dots, \mathbf{s}^{(N)} \in \mathcal{S}_N$ locate densely to each other so that their neighborhoods $\mathcal{N}^{(n)}$ (with radius $r^{(n)}$) overlap. Intuitively, if each $\mathcal{N}^{(n)}$ has a unified permutation for all $\mathbf{s}$ within $\mathcal{N}^{(n)}$, such overlapping leads to the same permutation for the union of all $\mathcal{N}^{(n)}$'s. Condition 5 means that there are enough points $\mathbf{s}^{(1)}, \dots, \mathbf{s}^{(n)}$ and they are densely positioned in $\mathcal{S}$, such that $\mathcal{N}$ covers a sufficiently large fraction of $\mathcal{S}$. This required fraction depends on the bi-Lipschitz constants with respect to Θ of Jacobian volume surrogate $V(\mathbf{s}; \Theta)$.

4 Related Works

IMA, Sparse Jacobian, and Object-Centric Learning. As mentioned in “Motivation” (see Sec. 3.1), IMA [20] and sparse Jacobian-based methods (e.g., [2, 19, 49, 50]) are conceptually related to our work. Additional remarks on connections between the SDI condition and IMA, Jacobian sparsity, anchor features [49], and object-centric methods [2, 23, 24], can be found in Appendix D.

SSC and Matrix Factorization. Volume-regularization approaches are popular in SMF models—i.e., $\mathbf{X} = \mathbf{W}\mathbf{H}$ with structural constraints (e.g., boundedness and nonnegativity) on $\mathbf{W}$ and/or $\mathbf{H}$ [27, 30–33]—where SSC-like conditions are imposed on the matrix factors to ensure identifiability of $\mathbf{W}$ and $\mathbf{H}$. The work [25] was the first to employ an SSC defined over the nonnegative orthant (originated from [26]) to show that latent-factor volume minimization identifies the corresponding SMF model. There, SSC requires that the columns/rows of a factor matrix widely spread in the nonnegative orthant. This line of work later was generalized to SSC in different norm balls [30, 31]. The SDI condition can be viewed as an extension of SSC in [30] (also see a similar SSC in [31, 33]) into the Jacobian domain over a continuous manifold. Although NMMI and SMF contexts differ substantially, some mathematical properties resulted from SSC (coupled with a volume-maximizing criterion) in [30] provide key stepping stones for identifiability analysis in this work.

Maximum Likelihood Estimation (MLE) and InfoMax for ICA. Both InfoMax [56] and MLE [57] handle ICA via optimizing the Jacobian volume in their formulations. Nonetheless, the Jacobian volume arises in their learning criteria as a result of information-theoretic or statistical estimation-based derivation (e.g., being surrogate of entropy). Model identifiability under nonlinear mixture models has not been established under these frameworks yet.

5 Numerical Results

Implementation¹. Given L realizations of $\mathbf{x}$, i.e., $\{\mathbf{x}^{(1)}, \mathbf{x}^{(2)}, \dots, \mathbf{x}^{(L)}\}$, we use multi-layer perceptrons (MLPs) to parametrize $\mathbf{f}_{\theta}$ and $\mathbf{g}_{\phi}$. We implement a regularized version of J-VolMax in (7) and train for T iterations via a warm-up heuristic with the number of warm-up epochs is $T_w < T$. Our

¹Our implementation code is available here: <https://github.com/hsnguyen24/dica>

loss function $\mathcal{L}_t$ at the t th epoch is

$$\min_{\theta, \phi} \mathcal{L}_t := \frac{1}{L} \sum_{n=1}^L \left(\|\mathbf{x}^{(n)} - \mathbf{f}_{\theta}(\mathbf{g}_{\phi}(\mathbf{x}^{(n)}))\|_2^2 - \lambda_{\text{vol}}(t) \times c_{\text{vol}} + \lambda_{\text{norm}} \times c_{\text{norm}}(t) \right). \quad (13)$$

In the above, c_{vol} is defined as

$$c_{\text{vol}} := \log \det(\mathbf{J}_{\mathbf{f}_{\theta}}(\mathbf{g}_{\phi}(\mathbf{x}^{(n)}))^{\top} \mathbf{J}_{\mathbf{f}_{\theta}}(\mathbf{g}_{\phi}(\mathbf{x}^{(n)}))), \quad (14)$$

and $\lambda_{\text{vol}}(t) := \frac{\lambda_{\text{vol}}}{T_w} \min\{t, T_w\}$; i.e., $\lambda_{\text{vol}}(t)$ linearly increases from $t = 0$ to a chosen $\lambda_{\text{vol}} > 0$ at the end of warm-up phase, i.e., $t = T_w$. Note that the explicit form of (14) might impose computational challenges when m and d are large. Hence, for high-dimensional datasets, one can use a trace-based surrogate of c_{vol} , reducing the per-iteration flops from $\mathcal{O}(m^3)$ to $\mathcal{O}(m^2)$; see Appendix E.1.

The term $c_{\text{norm}}(t)$ implements the norm constraint in (7b) via

$$c_{\text{norm}}(t) := \begin{cases} \|\mathbf{J}_{\mathbf{f}_{\theta}}(\mathbf{g}_{\phi}(\mathbf{x}^{(n)}))\|_1 & \text{if } t \leq T_w \\ \text{Softplus}\{\|\mathbf{J}_{\mathbf{f}_{\theta}}(\mathbf{g}_{\phi}(\mathbf{x}^{(n)}))\|_1 - C\} & \text{if } t > T_w \end{cases}, \quad (15)$$

with $\lambda_{\text{norm}} > 0$. The element-wise function $\text{Softplus}(z) = \ln(1 + e^z) \in \mathbb{R}$ is used as a smooth approximation of the hinge loss (see [58]). This term serves as a soft version of the norm constraint (7b) in J-VolMax; that is, it penalizes large L_1 norm values but does not put a hard constraint on their upper bound; see [59].

After the warm-up period, C is set to be the average of $\|\mathbf{J}_{\mathbf{f}_{\theta}}(\mathbf{g}_{\phi}(\mathbf{x}^{(n)}))\|_1$ in the last 10 epochs, to avoid that the c_{norm} term disproportionately skews the optimization process. More discussions on implementation can be found in Appendix E.1.

Evaluation. Following standard practice in NMMI [8, 12], we calculate both the mean Pearson correlation coefficient (MCC) and the mean coefficient of determination R^2 over pairs of d latent components, after fixing the permutation ambiguity via the Hungarian algorithm [60]. While MCC can only measure linear correlation, the R^2 score can measure nonlinear dependence. A perfect R^2 score means there is a mapping between the estimated latent component and the ground truth.

Baselines. We use the unregularized autoencoder (Base), autoencoder with the IMA contrast [20, 52] (IMA), and the sparsity-regularized loss (Sparse) as in (4). The L_1 -norm regularization with weight $\lambda_{\text{sp}} > 0$ is used to promote Jacobian sparsity, as in [19, 50]. To present the best performance of all the baselines under the considered settings, we tune their hyperparameters in a separately generated validation dataset with known ground-truth latent components. The hyperparameters λ_{vol} , λ_{norm} , and λ_{sp} are chosen by grid search from $\{10^{-2}, 10^{-3}, 10^{-4}, 10^{-5}\}$ on the same validation set.

5.1 Synthetic Data Experiments

Data Generation. We generate three types of mixtures to test our method. The generating process and ground-truth latent variables are unknown to our algorithms and the baselines. We use such controlled generation for constructing SDI-satisfying $\mathbf{f}$'s. In all simulations, we use two fully-connected ReLU neural networks with one hidden layer of 64 neurons to represent $\mathbf{f}_{\theta}$ and $\mathbf{g}_{\phi}$.

Mixture A: We use a linear mixture as the first checkpoint of our theory. Here, the latent components $\mathbf{s}$ are sampled from a normal distribution $p(\mathbf{s}) = \mathcal{N}(\mathbf{0}, \mathbf{\Sigma})$, where the covariance $\mathbf{\Sigma}$ is drawn from a Wishart distribution $W_p(\mathbf{I}, d)$. This way, the components of $\mathbf{s}$ are dependent. The observed mixtures are created by generating a matrix $\mathbf{A} \in \mathbb{R}^{m \times d}$ to form $\mathbf{x} = \mathbf{A}\mathbf{s}$. The matrix $\mathbf{A}$ can be generated to approximately satisfy the SDI condition by sampling m points in $\mathbb{R}^d$ randomly from an inflated L_2 ball, and then projecting those points onto the chosen weighted norm ball $\mathcal{B}_1^{w(s)}$ to create the m rows of $\mathbf{A}$; see [30] for a similar way to generate matrices that approximate SSC.

Mixture B: On top of $\mathbf{z} = \mathbf{A}\mathbf{s}$ as in Mixture A, we also apply an element-wise nonlinear distortion to create nonlinear mixtures, i.e., $\mathbf{x} = \mathbf{f}(\mathbf{z})$ and $[\mathbf{f}(\mathbf{z})]_i = a \cos(z_i) + z_i$, where $a \sim U(0.5, 1.0)$. Notice that $\mathbf{J}_{\mathbf{f}}(\mathbf{s}) = \text{Diag}(1 - a \sin(\mathbf{A}\mathbf{s}))\mathbf{A}$ under this construction. Therefore, with $\mathbf{A}$ approximately satisfying SDI and the nonlinear distortion weight a being sufficiently small, $\mathbf{J}_{\mathbf{f}}(\mathbf{s})$ also approximately satisfies the SDI.

Mixture C: This mixture is generated by $\mathbf{x} = \mathbf{f}(\mathbf{s})$, where $\mathbf{s} \sim U(-1, 1)$ and the nonlinear mixing function is $\mathbf{f}(\cdot) = (f_1(\cdot), \dots, f_m(\cdot)) \in \mathbb{R}^m$. The functions $f_k : \mathbb{R}^d \rightarrow \mathbb{R}, k = 1, \dots, m$ are MLPs

Table 1: Nonlinear R^2 and MCC scores (mean $\pm$ std., over 10 random trials).

Model	(d, m)	Mixture A		Mixture B		Mixture C	
		R^2	MCC	R^2	MCC	R^2	MCC
DICA	(2, 30)	0.92 ± 0.16	0.93 ± 0.15	0.97 ± 0.06	0.98 ± 0.03	0.95 ± 0.13	0.97 ± 0.09
	(3, 40)	0.92 ± 0.11	0.94 ± 0.08	0.99 ± 0.02	0.99 ± 0.01	0.90 ± 0.12	0.95 ± 0.07
	(4, 50)	0.98 ± 0.06	0.98 ± 0.04	0.92 ± 0.10	0.95 ± 0.07	0.87 ± 0.12	0.92 ± 0.07
	(5, 60)	0.92 ± 0.10	0.93 ± 0.09	0.89 ± 0.13	0.93 ± 0.08	0.87 ± 0.11	0.93 ± 0.06
Sparse	(2, 30)	0.79 ± 0.19	0.83 ± 0.12	0.88 ± 0.18	0.82 ± 0.12	0.87 ± 0.12	0.92 ± 0.07
	(3, 40)	0.71 ± 0.15	0.78 ± 0.13	0.72 ± 0.11	0.80 ± 0.09	0.71 ± 0.11	0.83 ± 0.07
	(4, 50)	0.65 ± 0.13	0.72 ± 0.11	0.71 ± 0.10	0.81 ± 0.06	0.64 ± 0.12	0.79 ± 0.07
	(5, 60)	0.63 ± 0.10	0.71 ± 0.07	0.67 ± 0.07	0.78 ± 0.05	0.60 ± 0.09	0.76 ± 0.06
Base	(2, 30)	0.88 ± 0.13	0.86 ± 0.11	0.81 ± 0.17	0.88 ± 0.11	0.77 ± 0.12	0.87 ± 0.08
	(3, 40)	0.69 ± 0.08	0.74 ± 0.07	0.68 ± 0.12	0.80 ± 0.09	0.62 ± 0.14	0.77 ± 0.09
	(4, 50)	0.54 ± 0.12	0.63 ± 0.09	0.63 ± 0.15	0.76 ± 0.10	0.48 ± 0.08	0.68 ± 0.06
	(5, 60)	0.42 ± 0.09	0.53 ± 0.07	0.54 ± 0.06	0.71 ± 0.04	0.47 ± 0.06	0.66 ± 0.04
IMA	(2, 30)	0.86 ± 0.14	0.92 ± 0.10	0.84 ± 0.13	0.91 ± 0.07	0.84 ± 0.14	0.91 ± 0.07
	(3, 40)	0.70 ± 0.10	0.83 ± 0.06	0.69 ± 0.13	0.82 ± 0.09	0.67 ± 0.07	0.81 ± 0.04
	(4, 50)	0.70 ± 0.10	0.83 ± 0.07	0.68 ± 0.09	0.81 ± 0.06	0.56 ± 0.11	0.74 ± 0.07
	(5, 60)	0.66 ± 0.05	0.80 ± 0.04	0.63 ± 0.10	0.79 ± 0.06	0.53 ± 0.08	0.72 ± 0.05

Table 2: Nonlinear R^2 scores for different ratios m/d (mean $\pm$ std., over 10 random trials).

(d, m)	DICA	IMA	Sparse	Base
(3, 40)	0.90 ± 0.10	0.63 ± 0.15	0.80 ± 0.13	0.63 ± 0.10
(3, 30)	0.89 ± 0.07	0.63 ± 0.12	0.77 ± 0.12	0.63 ± 0.14
(3, 20)	0.82 ± 0.17	0.60 ± 0.12	0.74 ± 0.13	0.60 ± 0.10
(3, 10)	0.71 ± 0.17	0.56 ± 0.09	0.75 ± 0.12	0.63 ± 0.16
(3, 5)	0.64 ± 0.09	0.66 ± 0.16	0.78 ± 0.16	0.61 ± 0.08
(3, 3)	0.55 ± 0.07	0.60 ± 0.09	0.58 ± 0.16	0.51 ± 0.09

whose layer weights are randomly generated. In addition, we pick $f_1, f_2, \dots, f_{\lfloor m/2 \rfloor}$ and modify them as follows: *a)* randomly choose $d - 1$ elements from $[d]$; and *b)* add an input layer that down-scales the chosen $d - 1$ latent component s_j by $\tilde{s}_j = \alpha_j \beta_j s_j$, where $\alpha_j \sim U(0.001, 0.002)$ and $\mathbb{P}(\beta_j = 1) = \mathbb{P}(\beta_j = -1) = 1/2$. This way, the first half of $\nabla f_k(s)$'s create a subset of gradients that are dominated by one latent component (cf. "Physical Meaning of SDI" in Section 3.1).

We should mention that, unlike Mixtures A and B, we have less control on generating SDI-satisfying mixtures under Mixture C. Hence, this type of nonlinear mixture can be used to test the model robustness of our approach.

Results. Table 1 shows both the MCC and the nonlinear R^2 scores of all four methods, namely, DICA, IMA, Sparse, and Base, on Mixtures A, B, C. We observe that for Mixtures A and B, both scores attained by DICA are mostly above 0.92 (except for R^2 under $(d, m) = (5, 60)$). In contrast, IMA, Sparse, and Base do not perform as well. Importantly, the performance of DICA does not significantly degrade as m and d grows, but the baselines appear to struggle under larger m 's and d 's. For the more challenging Mixture C, DICA's performance scores still have a substantial margin over the baselines, and the scores remain largely the same when (d, m) varies.

Ablation on m/d . We examine the performance of DICA when varying the ratio m/d . Per our discussion (cf. Sec. 3), SDI favors $m \gg d$ cases. As the ratio m/d approaches 2, SDI is less likely to be satisfied. In Table 2, we use Mixture C to test the performance of DICA where we fix $d = 3$ and vary m . Consistent with our analysis, the performance of DICA deteriorates as m approaches $2d$. For $m < 2d$, the performance of DICA is similar to vanilla autoencoder—showing that the effect of volume regularization vanishes in these cases. This result confirms that DICA is suitable for learning low-dimensional latent components from high-dimensional data, e.g., image representation learning.

5.2 Inferring Transcription Factor Activities from Single-cell Gene Expressions

In this experiment, we apply J-VolMax (7) to inferring transcription factors (TFs) activities in single-cell transcriptomics analysis [61]. Specifically, we employ J-VolMax to infer the ground-truth mRNA concentrations of TFs from gene expression data. In this context, our single-cell gene expressions are from a bio-realistic generator, namely, SERGIO [62]; see the use of SERGIO in the NMMI literature [41, 63]. The mRNA concentration levels $s \in \mathbb{R}^d$ of the TFs are governed from a chemical Langevin

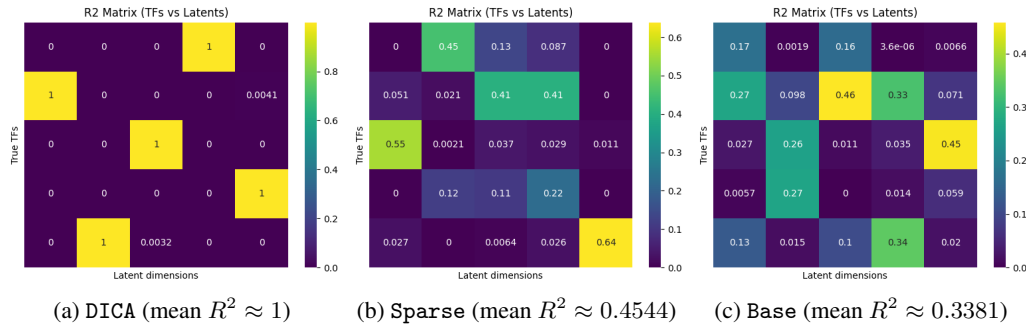


Figure 2: Heatmap of R^2 scores between estimated components and ground-truth mRNA concentrations of TFs.

equation, and the gene expressions (i.e., samples of $\mathbf{x} \in \mathbb{R}^m$) are generated by a gene regulatory mechanism $\mathbf{f} : \mathbb{R}^d \mapsto \mathbb{R}^m$ applied onto $\mathbf{s}$.

We use the TRRUST dataset of TF-gene pairs of mouse [64]. It is a manually curated dataset that includes high-confidence interactions between TFs and their target genes. These interactions are either activating or repressing. All entries are supported by experimental evidence from multiple published studies. We sample a sub-network from the large gene regulatory network provided by TRRUST, via which $m = 178$ genes are governed by $d = 5$ master regulator TFs. More details of the data generation process using SERGIO is in Appendix E.3.

Fig. 2 shows average R^2 scores (from 10 trials) for each $s_1, \dots, s_5$ up to best permutation, as obtained by DICA, Sparse, and Base. One can see that DICA successfully infers the ground-truth mRNA concentration level $s_1, \dots, s_5$ of each TF (up to a permutation and invertible mapping), but Sparse and Base are not able to attain reasonable R^2 scores. Furthermore, the proposed method can disentangle the influences of different TFs. Specifically, each learned latent component is only matched to a unique ground-truth TF with R^2 score ≈ 1 , yet the R^2 score is approximately 0 when matching with other TFs. In contrast, Base could not disentangle the TFs at all. Similarly, Sparse could only mildly disentangle two TFs (i.e., the 3th and the 5th TFs). Additionally, the MCC score of DICA reaches ≈ 1 , whereas Sparse and Base output $MCC \approx 0.6635$ and $MCC \approx 0.5721$, respectively.

Some remarks regarding this experiment is as follows. Note that the data generation process follows biology-driven mechanisms in SERGIO, and thus we do not have a way to enforce SDI. However, as the TFs are believed to have diverse influences on the gene expressions (e.g., due to the activating/repressing effects on genes by the TFs), it is reasonable to assume that SDI approximately holds in this application—which perhaps explains why DICA works well. On the other hand, some TFs may influence many gene expressions ([65]), and many cross-interactions between TFs and genes exist ([66]). Hence, the assumption that the Jacobian of $\mathbf{f}$ is strictly sparse might be stringent. Therefore, the sparse Jacobian approach may not be a good fit, as its performance indicates.

6 Conclusion

We introduced DICA, a new paradigm for nonlinear mixture learning with identifiability guarantees. Via the insight of convex geometry, we showed that, as long as the latent components have sufficiently different influences on the observed features so that a geometric condition, namely, the SDI condition, is satisfied, the mixture model is identifiable via fitting the generative model and maximizing the Jacobian of the learning function simultaneously. This J-VolMax criterion ensures identifiability without relying on many assumptions in the NMML literature, e.g., availability of auxiliary information, independence of latent components, or the sparsity of the mixing function's Jacobian, thus offering a valuable alternative to existing NMML methods. The experiments supported our theory.

Limitations and Future Works. We noticed a number of limitations. First, the DICA framework has a log det term to compute, which poses a challenging optimization problem. This is particularly hard as this term has to be evaluated at every realization of $\mathbf{s}$. How to optimize the J-VolMax criterion efficiently is an interesting yet challenging direction to consider. Second, the identifiability analysis was done under ideal conditions where noise is absent, one has access to unlimited samples, and the learner class is assumed to perfectly include the target function. While J-VolMax is empirically shown to be quite robust in experiments, a more thorough identifiability analysis under practical conditions would close the gap, e.g., by analyzing sample complexity under noise and model mismatches.

Acknowledgment: This work is supported in part by the National Science Foundation (NSF) CAREER Award ECCS-2144889. The authors would like to thank Sagar Shrestha for insightful discussions on the identifiability analysis, and the anonymous reviewers for constructive feedbacks.

References

- [1] Ilyes Khemakhem et al. “Variational Autoencoders and Nonlinear ICA: A Unifying Framework”. In: *Proceedings of the Twenty Third International Conference on Artificial Intelligence and Statistics*. Vol. 108. Proceedings of Machine Learning Research. PMLR, 2020, pp. 2207–2217.
- [2] Jack Brady et al. “Provably Learning Object-Centric Representations”. In: *Proceedings of the 40th International Conference on Machine Learning*. Vol. 202. Proceedings of Machine Learning Research. PMLR, 2023, pp. 3038–3062.
- [3] Yujia Zheng and Kun Zhang. “Generalizing Nonlinear ICA Beyond Structural Sparsity”. In: *Advances in Neural Information Processing Systems*. Vol. 36. 2023, pp. 13326–13355.
- [4] Irina Higgins et al. “beta-VAE: Learning basic visual concepts with a constrained variational framework.” In: *International Conference on Learning Representations* (2017).
- [5] Bernhard Schölkopf et al. “Toward Causal Representation Learning”. In: *Proceedings of the IEEE* 109.5 (2021), pp. 612–634.
- [6] Julius von Kügelgen. “Identifiable Causal Representation Learning: Unsupervised, Multi-View, and Multi-Environment”. PhD thesis. University of Cambridge, 2023.
- [7] Qi Lyu et al. “Understanding Latent Correlation-Based Multiview Learning and Self-Supervision: An Identifiability Perspective”. In: *International Conference on Learning Representations*. 2022.
- [8] Julius von Kügelgen et al. “Self-Supervised Learning with Data Augmentations Provably Isolates Content from Style”. In: *Advances in Neural Information Processing Systems*. Vol. 34. 2021, pp. 16451–16467.
- [9] Aapo Hyvärinen and Petteri Pajunen. “Nonlinear independent component analysis: Existence and uniqueness results”. In: *Neural Networks* 12.3 (1999), pp. 429–439.
- [10] Aapo Hyvarinen and Hiroshi Morioka. “Unsupervised Feature Extraction by Time-Contrastive Learning and Nonlinear ICA”. In: *Advances in Neural Information Processing Systems*. Vol. 29. 2016.
- [11] Aapo Hyvarinen and Hiroshi Morioka. “Nonlinear ICA of Temporally Dependent Stationary Sources”. In: *Proceedings of the 20th International Conference on Artificial Intelligence and Statistics*. Vol. 54. Proceedings of Machine Learning Research. PMLR, 2017, pp. 460–469.
- [12] Aapo Hyvarinen, Hiroaki Sasaki, and Richard Turner. “Nonlinear ICA Using Auxiliary Variables and Generalized Contrastive Learning”. In: *Proceedings of the Twenty-Second International Conference on Artificial Intelligence and Statistics*. Vol. 89. Proceedings of Machine Learning Research. PMLR, 2019, pp. 859–868.
- [13] S. Achard and C. Jutten. “Identifiability of Post-Nonlinear Mixtures”. In: *IEEE Signal Processing Letters* 12.5 (2005), pp. 423–426.
- [14] Qi Lyu and Xiao Fu. “Provable Subspace Identification Under Post-Nonlinear Mixtures”. In: *Advances in Neural Information Processing Systems*. Vol. 35. 2022, pp. 1554–1567.
- [15] Qi Lyu and Xiao Fu. “Identifiability-Guaranteed Simplex-Structured Post-Nonlinear Mixture Learning via Autoencoder”. In: *IEEE Transactions on Signal Processing* 69 (2021).
- [16] Andreas Ziehe et al. “Blind Separation of Post-nonlinear Mixtures using Linearizing Transformations and Temporal Decorrelation”. In: *Journal of Machine Learning Research* (2003), pp. 1319–1338.
- [17] A. Hyvarinen and P. Pajunen. “On existence and uniqueness of solutions in nonlinear independent component analysis”. In: *1998 IEEE International Joint Conference on Neural Networks Proceedings. IEEE World Congress on Computational Intelligence*. Vol. 2. 1998, pp. 1350–1355.
- [18] Bohdan Kivva et al. “Identifiability of deep generative models without auxiliary information”. In: *Advances in Neural Information Processing Systems*. Vol. 35. 2022, pp. 15687–15701.

- [19] Yujia Zheng, Ignavier Ng, and Kun Zhang. “On the Identifiability of Nonlinear ICA: Sparsity and Beyond”. In: *Advances in Neural Information Processing Systems*. Vol. 35. 2022, pp. 16411–16422.
- [20] Luigi Gresele et al. “Independent mechanism analysis, a new concept?” In: *Advances in Neural Information Processing Systems*. Vol. 34. 2021, pp. 28233–28248.
- [21] Sebastien Lachapelle et al. “Disentanglement via Mechanism Sparsity Regularization: A New Principle for Nonlinear ICA”. In: *Proceedings of the First Conference on Causal Learning and Reasoning*. Vol. 177. Proceedings of Machine Learning Research. PMLR, 2022, pp. 428–484.
- [22] Bernhard Schölkopf, Jonas Peters, and Dominik Janzing. *Elements of Causal Inference*. Adaptive Computation and Machine Learning series. MIT Press, 2017.
- [23] Sébastien Lachapelle et al. “Additive Decoders for Latent Variables Identification and Cartesian-Product Extrapolation”. In: *Advances in Neural Information Processing Systems*. Vol. 36. 2023, pp. 25112–25150.
- [24] Jack Brady et al. “Interaction Asymmetry: A General Principle for Learning Composable Abstractions”. In: *The Thirteenth International Conference on Learning Representations*. 2025.
- [25] Xiao Fu et al. “Blind Separation of Quasi-Stationary Sources: Exploiting Convex Geometry in Covariance Domain”. In: *IEEE Transactions on Signal Processing* 63.9 (2015), pp. 2306–2320.
- [26] Kejun Huang, Nicholas D. Sidiropoulos, and Ananthram Swami. “Non-Negative Matrix Factorization Revisited: Uniqueness and Algorithm for Symmetric Decomposition”. In: *IEEE Transactions on Signal Processing* 62.1 (2014), pp. 211–224.
- [27] Xiao Fu, Kejun Huang, and Nicholas D. Sidiropoulos. “On Identifiability of Nonnegative Matrix Factorization”. In: *IEEE Signal Processing Letters* 25.3 (2018), pp. 328–332.
- [28] Nicolas Gillis. *Nonnegative Matrix Factorization*. Society for Industrial and Applied Mathematics, Jan. 2020. ISBN: 9781611976410.
- [29] Xiao Fu et al. “Nonnegative Matrix Factorization for Signal and Data Analytics: Identifiability, Algorithms, and Applications”. In: *IEEE Signal Processing Magazine* 36.2 (2019), pp. 59–80.
- [30] Gokcan Tatli and Alper T. Erdogan. “Polytopic Matrix Factorization: Determinant Maximization Based Criterion and Identifiability”. In: *IEEE Transactions on Signal Processing* 69 (2021), pp. 5431–5447.
- [31] Jingzhou Hu and Kejun Huang. “Global Identifiability of L1-based Dictionary Learning via Matrix Volume Optimization”. In: *Advances in Neural Information Processing Systems*. Vol. 36. 2023, pp. 36165–36186.
- [32] Jingzhou Hu and Kejun Huang. “Identifiable Bounded Component Analysis Via Minimum Volume Enclosing Parallelootope”. In: *2023 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*. 2023.
- [33] Yuchen Sun and Kejun Huang. “Global Identifiability of Overcomplete Dictionary Learning via L1 and Volume Minimization”. In: *The Thirteenth International Conference on Learning Representations*. 2025.
- [34] Xiao Fu et al. “Robust Volume Minimization-Based Matrix Factorization for Remote Sensing and Document Clustering”. In: *IEEE Transactions on Signal Processing* 64.23 (2016), pp. 6254–6268.
- [35] Lidan Miao and Hairong Qi. “Endmember extraction from highly mixed data using minimum volume constrained nonnegative matrix factorization”. In: *IEEE Transactions on Geoscience and Remote Sensing* 45.3 (2007), pp. 765–777.
- [36] Chia-Hsiang Lin et al. “Maximum Volume Inscribed Ellipsoid: A New Simplex-Structured Matrix Factorization Framework via Facet Enumeration and Convex Optimization”. In: *SIAM Journal on Imaging Sciences* 11.2 (2018), pp. 1651–1679.
- [37] Christian Jutten and Jeanny Herault. “Blind separation of sources, part I: An adaptive algorithm based on neuromimetic architecture”. In: *Signal Processing* 24.1 (1991), pp. 1–10.
- [38] Pierre Comon. “Independent component analysis, A new concept?” In: *Signal Processing* 36.3 (1994), pp. 287–314. ISSN: 0165-1684.
- [39] Sébastien Lachapelle et al. *Nonparametric Partial Disentanglement via Mechanism Sparsity: Sparse Actions, Interventions and Sparse Temporal Dependencies*. 2024. arXiv: 2401.04890.
- [40] David A. Klindt et al. “Towards Nonlinear Disentanglement in Natural Data with Temporal Sparse Coding”. In: *International Conference on Learning Representations*. 2021.

- [41] Hiroshi Morioka and Aapo Hyvarinen. “Causal Representation Learning Made Identifiable by Grouping of Observational Variables”. In: *Proceedings of the 41st International Conference on Machine Learning*. Vol. 235. Proceedings of Machine Learning Research. PMLR, 2024, pp. 36249–36293.
- [42] Luigi Gresele et al. “The Incomplete Rosetta Stone problem: Identifiability results for Multi-view Nonlinear ICA”. In: *Proceedings of The 35th Uncertainty in Artificial Intelligence Conference*. Proceedings of Machine Learning Research. PMLR, 2020.
- [43] Sagar Shrestha and Xiao Fu. “Towards Identifiable Unsupervised Domain Translation: A Diversified Distribution Matching Approach”. In: *The Twelfth International Conference on Learning Representations*. 2024.
- [44] Simon Buchholz, Michel Besserve, and Bernhard Schölkopf. “Function Classes for Identifiable Nonlinear Independent Component Analysis”. In: *Advances in Neural Information Processing Systems*. Vol. 35. 2022, pp. 16946–16961.
- [45] Daniella Horan, Eitan Richardson, and Yair Weiss. “When Is Unsupervised Disentanglement Possible?” In: *Advances in Neural Information Processing Systems*. 2021, pp. 5150–5161.
- [46] Kun Zhang and Laiwan Chan. “Minimal Nonlinear Distortion Principle for Nonlinear Independent Component Analysis”. In: *Journal of Machine Learning Research* 9.81 (2008), pp. 2455–2487.
- [47] A. Taleb and C. Jutten. “Source Separation in Post-nonlinear Mixtures”. In: *IEEE Transactions on Signal Processing* 47.10 (1999), pp. 2807–2820.
- [48] Avinash Kori et al. “Identifiable Object-Centric Representation Learning via Probabilistic Slot Attention”. In: *Advances in Neural Information Processing Systems*. Vol. 37. 2024, pp. 93300–93335.
- [49] Gemma Elyse Moran et al. “Identifiable Deep Generative Models via Sparse Decoding”. In: *Transactions on Machine Learning Research* (2022). ISSN: 2835-8856.
- [50] Travers Rhodes and Daniel Lee. “Local Disentanglement in Variational Auto-Encoders Using Jacobian L_1 Regularization”. In: *Advances in Neural Information Processing Systems*. Vol. 34. 2021, pp. 22708–22719.
- [51] William Peebles et al. “The Hessian Penalty: A Weak Prior for Unsupervised Disentanglement”. In: *Proceedings of European Conference on Computer Vision (ECCV)*. 2020.
- [52] Shubhangi Ghosh et al. “Independent Mechanism Analysis and the Manifold Hypothesis”. In: *Causal Representation Learning Workshop at NeurIPS 2023*. 2023.
- [53] Francesco Locatello et al. “Challenging Common Assumptions in the Unsupervised Learning of Disentangled Representations”. In: *Proceedings of the 36th International Conference on Machine Learning*. Proceedings of Machine Learning Research. PMLR, 2019, pp. 4114–4124.
- [54] Shahana Ibrahim et al. “Crowdsourcing via Pairwise Co-occurrences: Identifiability and Algorithms”. In: *Advances in Neural Information Processing Systems*. 2019.
- [55] Peter L Bartlett and Shahar Mendelson. “Rademacher and gaussian complexities: Risk bounds and structural results”. In: *Journal of machine learning research* 3.Nov (2002), pp. 463–482.
- [56] Anthony J. Bell and Terrence J. Sejnowski. “An Information-Maximization Approach to Blind Separation and Blind Deconvolution”. In: *Neural Computation* 7.6 (1995), pp. 1129–1159. ISSN: 1530-888X.
- [57] J.-F. Cardoso. “Blind signal separation: statistical principles”. In: *Proceedings of the IEEE* 86.10 (1998), pp. 2009–2025.
- [58] Xavier Glorot, Antoine Bordes, and Yoshua Bengio. “Deep Sparse Rectifier Neural Networks”. In: *Proceedings of the Fourteenth International Conference on Artificial Intelligence and Statistics*. Vol. 15. Proceedings of Machine Learning Research. PMLR, 2011, pp. 315–323.
- [59] Jose M. Bioucas-Dias. “A variable splitting augmented Lagrangian approach to linear spectral unmixing”. In: *2009 First Workshop on Hyperspectral Image and Signal Processing: Evolution in Remote Sensing*. 2009.
- [60] H. W. Kuhn. “The Hungarian method for the assignment problem”. In: *Naval Research Logistics Quarterly* 2.1–2 (Mar. 1955), pp. 83–97.
- [61] Dennis Hecker et al. “Computational tools for inferring transcription factor activity”. In: *Proteomics* 23 (2023). ISSN: 1615-9861.
- [62] Payam Dibaeinia and Saurabh Sinha. “SERGIO: A Single-Cell Expression Simulator Guided by Gene Regulatory Networks”. In: *Cell Systems* 11.3 (2020), 252–271.e11. ISSN: 2405-4712.

- [63] Hiroshi Morioka and Aapo Hyvarinen. “Connectivity-contrastive learning: Combining causal discovery and representation learning for multimodal data”. In: *Proceedings of The 26th International Conference on Artificial Intelligence and Statistics*. Vol. 206. Proceedings of Machine Learning Research. PMLR, 2023, pp. 3399–3426.
- [64] Heonjong Han et al. “TRRUST v2: an expanded reference database of human and mouse transcriptional regulatory interactions”. In: *Nucleic Acids Research* 46.D1 (2017).
- [65] Samuel A. Lambert et al. “The Human Transcription Factors”. In: *Cell* 172.4 (2018), pp. 650–665. ISSN: 0092-8674.
- [66] Tamar Friedlander et al. “Intrinsic limits to gene regulation by global crosstalk”. In: *Nature Communications* 7 (2016). ISSN: 2041-1723.
- [67] Stephen Boyd and Lieven Vandenbergh. *Convex Optimization*. Cambridge University Press, 2004.
- [68] Arne Brøndsted. *An Introduction to Convex Polytopes*. Graduate Texts in Mathematics. Springer, 2012.
- [69] Loring W Tu. *An Introduction to Manifolds*. 2nd ed. Universitext. Springer, 2010.
- [70] Shai Shalev-Shwartz and Shai Ben-David. *Understanding machine learning*. Cambridge University Press, 2014.
- [71] Yuxiang Wei et al. “Orthogonal Jacobian Regularization for Unsupervised Disentanglement in Image Generation”. In: *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)*. 2021.
- [72] M. Berman et al. “ICE: a statistical approach to identifying endmembers in hyperspectral images”. In: *IEEE Transactions on Geoscience and Remote Sensing* (2004), pp. 2085–2095.
- [73] Luigi Gresele et al. “Relative gradient optimization of the Jacobian term in unsupervised deep learning”. In: *Advances in Neural Information Processing Systems*. Vol. 33. 2020.
- [74] Kevin P. Murphy. *Probabilistic Machine Learning: An Introduction*. MIT Press, 2022.
- [75] Diederik P. Kingma and Jimmy Ba. “Adam: A Method for Stochastic Optimization”. In: *International Conference for Learning Representations*. 2015.
- [76] Kaiming He et al. “Delving Deep into Rectifiers: Surpassing Human-Level Performance on ImageNet Classification”. In: *2015 IEEE International Conference on Computer Vision (ICCV)*. 2015, pp. 1026–1034.
- [77] Y. LeCun et al. “Gradient-based learning applied to document recognition”. In: *Proceedings of the IEEE* 86.11 (1998), pp. 2278–2324.

Supplementary Material of “Diverse Influence Component Analysis: A Geometric Approach to Nonlinear Mixture Identifiability”

A Preliminaries

A.1 Notation

- For $m \in \mathbb{N}$, $[m]$ is a shorthand for $\{1, 2, \dots, m\}$.
- $x, \mathbf{x}$ and $\mathbf{X}$ denotes a scalar, a vector, and a matrix, respectively.
- $\mathbf{X}_{i,j}$, $[\mathbf{X}]_{i,j}$ and $x_{i,j}$ denotes the entry at i th row and j th column of $\mathbf{X}$. $\mathbf{X}_{i:}$ and $\mathbf{X}_{:j}$ denotes the i th row and j th column of $\mathbf{X}$.
- For a vector, $\|\cdot\|_p$ denotes L_p vector norm; for a matrix, $\|\cdot\|_p$ denotes the entry-wise L_p matrix norm and $\|\cdot\|$ denotes the spectral norm.
- The Jacobian of a function $\mathbf{f} : \mathbb{R}^d \rightarrow \mathbb{R}^m$ at point $\mathbf{s}$ is a matrix $\mathbf{J}_{\mathbf{f}}(\mathbf{s})$ of partial derivatives such that $[\mathbf{J}_{\mathbf{f}}(\mathbf{s})]_{i,j} = \partial f_i / \partial s_j$; $\nabla f(\mathbf{s}) = [\partial f(\mathbf{s}) / \partial s_1 \dots \partial f(\mathbf{s}) / \partial s_d]^\top$ denotes the gradient of scalar function $f : \mathbb{R}^d \rightarrow \mathbb{R}$; hence, $[\mathbf{J}_{\mathbf{f}}(\mathbf{s})]_{i:} = \nabla f_i(\mathbf{s})^\top$.
- $\text{conv}\{\cdot\}$ denotes the convex hull of a set of vectors.
- $\mathcal{B}_p(R)$ being the unweighted L_p -norm ball with radius $R > 0$, and $\mathcal{B}_p$ is shorthand for the unweighted unit L_p -norm ball (i.e., $R = 1$).
- $\mathcal{B}_1^{w(\mathbf{s})} = \{\mathbf{y} \in \mathbb{R}^d : \sum_{k=1}^d w_k(\mathbf{s})^{-1} |y_k| \leq 1\}$ is an L_1 -norm ball weighted by $1/w_1(\mathbf{s}), \dots, 1/w_d(\mathbf{s}) > 0$.
- $\mathcal{E}(P)$ is the maximum volume inscribed ellipsoid (MVIE) of an origin-centered convex polytope P .
- Polar of a set $\mathcal{C}$ is $\mathcal{C}^* = \{\mathbf{x} \in \mathbb{R}^d : \mathbf{x}^\top \mathbf{y} \leq 1, \forall \mathbf{y} \in \mathcal{C}\}$.
- $\text{bd}(P)$ is the boundary of a set P .
- $\text{extr}(P)$ is the set of extreme points (i.e., vertices) of a polytope P .
- $\mathcal{P}_d$ is the set of $d \times d$ permutation matrices.
- $\bar{\Theta} = [\bar{\theta}, \bar{\phi}]$ is shorthand for the vector of parameters of parametrized encoder $\mathbf{f}_{\bar{\theta}}$ and decoder $\mathbf{g}_{\bar{\phi}}$.
- $V(\mathbf{s}; \bar{\Theta}) := \log \det(\mathbf{J}_{\mathbf{f}_{\bar{\theta}}}(\mathbf{g}_{\bar{\phi}}(\mathbf{x}))^\top \mathbf{J}_{\mathbf{f}_{\bar{\theta}}}(\mathbf{g}_{\bar{\phi}}(\mathbf{x})))$ is shorthand for the log-determinant of Jacobian matrix of $\bar{\Theta} := [\bar{\theta}, \bar{\phi}]$, evaluated at $\mathbf{x} = \mathbf{f}(\mathbf{s})$; this is used as the Jacobian volume surrogate in J-VolMax objective (7a).

A.2 Background on Convex Geometry

In this section, we briefly introduce the key notions in convex geometry that will be used in the paper.

Convex Hull. Given a set $S \subset \mathbb{R}^d$, its convex hull $\text{conv}\{S\}$ is the set of all convex combinations of vectors in S , i.e.,

$$\text{conv}\{S\} = \left\{ \sum_{i=1}^m \alpha_i \mathbf{x}_i : \mathbf{x}_i \in S, \alpha_i \geq 0, \sum_{i=1}^m \alpha_i = 1 \right\}.$$

Maximum Volume Inscribed Ellipsoid (MVIE). An MVIE (sometimes called *John ellipsoid*) of a bounded convex set C is the ellipsoid of maximum volume that lies inside C , which can be represented as

$$\mathcal{E}(C) = \{\mathbf{B}\mathbf{u} + \mathbf{d} : \|\mathbf{u}\|_2 \leq 1\}$$

for $\mathbf{B} \in \mathbb{R}^{d \times d}$, $\mathbf{d} \in \mathbb{R}^d$, and $\mathbf{B}$ is a symmetric positive-definite matrix (i.e., $\mathbf{B} \in \mathbb{S}_{++}^d$). $\mathcal{E}(C)$ can be obtained from solving

$$\max_{\mathbf{B} \in \mathbb{S}_{++}^d, \mathbf{d} \in \mathbb{R}^d} \log \det \mathbf{B} \quad \text{s.t.} \quad \sup_{\mathbf{u} : \|\mathbf{u}\|_2 \leq 1} I_C(\mathbf{B}\mathbf{u} + \mathbf{d}) \leq 0,$$

with $I_C(\cdot)$ being the indicator function of whether a point is in C . Note that every non-empty bounded convex set has a unique MVIE. See [67, Section 8.4.2] for more details. Note that the MVIE has the following linear scaling property:

Proposition A.1. *If $T : \mathbb{R}^d \mapsto \mathbb{R}^d$ is an invertible linear map and C is a non-empty compact convex set, then the MVIE of $T(C)$ is the image under T of the MVIE of C , i.e.*

$$\mathcal{E}(T(C)) = T(\mathcal{E}(C)).$$

Proof. See [67, Section 8.4.3]. □

In the following, we give examples of MVIEs of several sets that are relevant to our context.

Proposition A.2. *The followings are true about a convex set and its MVIE:*

1. *The MVIE of an L_1 -norm ball $\mathcal{B}_1(R)$ in $\mathbb{R}^d$ is $\{\mathbf{x} : \|\mathbf{x}\|_2 \leq \frac{R}{\sqrt{d}}\}$;*
2. *The MVIE of an L_∞ -norm ball $\mathcal{B}_\infty(R)$ in $\mathbb{R}^d$ is $\{\mathbf{x} : \|\mathbf{x}\|_2 \leq R\}$;*
3. *The MVIE of an L_2 -norm ball $\mathcal{B}_2(R)$ is itself.*

Polytope. Intuitively, a polytope P is a bounded geometric object with flat sides (called *facets*). A polytope P containing the origin can be represented either by intersections of f half-spaces as

$$P = \{\mathbf{x} \in \mathbb{R}^d : \mathbf{a}_i^\top \mathbf{x} \leq 1, i = 1, \dots, f\},$$

or by convex hull of k vertices as

$$P = \text{conv}\{\mathbf{v}_1, \dots, \mathbf{v}_k\}.$$

For examples, L_p -norm ball $\mathcal{B}_p$ and simplex $\Delta = \{\mathbf{x} \in \mathbb{R}^d : x_i \geq 0, \sum_{i=1}^d x_i = 1\}$ are polytopes.

Polar Set. The polar of a set $C \subset \mathbb{R}^d$ is

$$C^* = \{\mathbf{x} \in \mathbb{R}^d : \mathbf{x}^\top \mathbf{y} \leq 1, \forall \mathbf{y} \in C\}.$$

There are some interesting properties of polar sets:

Proposition A.3. *For two sets $A, B \subset \mathbb{R}^d$ and $\alpha > 0$, the followings holds:*

1. *$A \subseteq B$ implies $B^* \subseteq A^*$,*
2. *$(A \cup B)^* = A^* \cap B^*$,*
3. *$(\alpha A)^* = \frac{1}{\alpha} A^*$.*

We give examples of polar sets of several sets that are related to our context:

Proposition A.4. *The followings are true about a set and its polar set:*

1. *Polar set of a weighted L_1 -norm ball $\mathcal{B}_1^w = \{\mathbf{y} \in \mathbb{R}^d : \sum_{k=1}^d w_k^{-1} |y_k| \leq 1\}$ is the reciprocally weighted L_∞ -norm ball $\mathcal{B}_\infty^{w^{-1}} = \{\mathbf{y} \in \mathbb{R}^d : \max_{k=1, \dots, d} \{w_k |y_k|\} \leq 1\}$, and vice versa;*
2. *Polar set of a L_2 -norm ball $\mathcal{B}_2(R)$ with radius $R > 0$ is a L_2 -norm ball $\mathcal{B}_2(\frac{1}{R})$ with radius $\frac{1}{R}$, and vice versa.*

For the details of Proposition A.3 and Proposition A.4, we refer the readers to [68, Chapter 2].

B Proof of Theorem 3.2

The proof consists of two parts. In the first part, i.e., Lemma B.2, we show that any optimal solution of J-VolMax problem (7) must identify a ground-truth latent vector $\mathbf{s} \in \mathcal{S}$, up to s -dependent permutation and invertible component-wise transformation. Then, in the second part, namely, Theorem 3.2, we leverage the continuity and full-rank structure of the Jacobian matrices at different $\mathbf{s}$ to argue that the permutation of latent components should be identical for all $\mathbf{s} \in \mathcal{S}$, instead of being s -dependent as initially shown in Lemma B.2. That gives us the desired identifiability result.

Before we begin the proof, we invoke the following lemma from [30]:

Lemma B.1. Let $\mathbf{H} \in \mathbb{R}^{d \times d}$ be a matrix satisfying

$$\|\mathbf{H}^\top \mathbf{u}\|_2 \leq \sqrt{d}, \forall \mathbf{u} \in \text{extr}(\mathcal{B}_\infty).$$

Then, $|\det \mathbf{H}| \leq 1$, with equality occurs if and only if $\mathbf{H}$ is a real orthogonal matrix.

Proof. See [30, Theorem 3]. □

Lemma B.2. Denote any optimal solution of Problem (7) as $(\hat{\boldsymbol{\theta}}, \hat{\boldsymbol{\phi}})$. Assume $\hat{\mathbf{f}} = \mathbf{f}_{\hat{\boldsymbol{\theta}}}$ and $\hat{\mathbf{g}} = \mathbf{g}_{\hat{\boldsymbol{\phi}}}$ are universal function representers. Suppose the model in (1) and Assumption 3.1 hold every point $\mathbf{s} \in \mathcal{S}$. Then, the Jacobian at the point $\mathbf{s} \in \mathcal{S}$ of the function $\mathbf{h}^* = \hat{\mathbf{g}} \circ \mathbf{f}$ will satisfy

$$\mathbf{J}_{\mathbf{h}^*}(\mathbf{s}) = \mathbf{D}(\mathbf{s})\boldsymbol{\Pi}(\mathbf{s}),$$

where $\mathbf{D}(\mathbf{s})$ is an invertible diagonal matrix with $[\mathbf{D}(\mathbf{s})]_{i,i}$ dependent on i -th component s_i only, and $\boldsymbol{\Pi}(\mathbf{s})$ is a permutation matrix dependent on the point $\mathbf{s}$, almost everywhere.

Proof. Define a differentiable function $\mathbf{h} : \mathcal{S} \mapsto \mathcal{S}$ as a composition $\mathbf{h} = \hat{\mathbf{g}} \circ \mathbf{f}$. Our ultimate goal is to show that the Jacobian $\mathbf{J}_{\mathbf{h}}(\mathbf{s})$ has a permutation-like support structure, which is equivalent to the result statement.

We begin our analysis by considering a data point $\mathbf{x} = \mathbf{f}(\mathbf{s})$ for $\mathbf{s} \in \mathcal{S}$. Let $\hat{\mathbf{s}} = \hat{\mathbf{g}}(\mathbf{x})$. By definition of function $\mathbf{h}$,

$$\mathbf{f}(\mathbf{s}) = \mathbf{x} = \mathbf{f}_{\hat{\boldsymbol{\theta}}}(\hat{\mathbf{s}}) = \mathbf{f}_{\hat{\boldsymbol{\theta}}}(\mathbf{g}_{\hat{\boldsymbol{\phi}}}(\mathbf{x})) = \mathbf{f}_{\hat{\boldsymbol{\theta}}}(\mathbf{h}(\mathbf{s})) \quad (16)$$

Then, by chain rule,

$$\mathbf{J}_{\mathbf{f}}(\mathbf{s}) = \mathbf{J}_{\mathbf{f}_{\hat{\boldsymbol{\theta}}}}(\mathbf{h}(\mathbf{s}))\mathbf{J}_{\mathbf{h}}(\mathbf{s}) = \mathbf{J}_{\mathbf{f}_{\hat{\boldsymbol{\theta}}}}(\hat{\mathbf{s}})\mathbf{J}_{\mathbf{h}}(\mathbf{s}). \quad (17)$$

We now argue that $\mathbf{J}_{\mathbf{h}}(\mathbf{s})$ is a full rank matrix. By definition of $\mathbf{h}$,

$$\mathbf{J}_{\mathbf{h}}(\mathbf{s}) = \mathbf{J}_{\mathbf{g}_{\hat{\boldsymbol{\phi}}}}(\mathbf{x})\mathbf{J}_{\mathbf{f}}(\mathbf{s}). \quad (18)$$

By definition of the ground-truth mixing process, $\mathbf{f}$ is an injective function, implying that $\mathbf{J}_{\mathbf{f}}(\mathbf{s})$ is full rank. We now investigate $\mathbf{J}_{\mathbf{g}_{\hat{\boldsymbol{\phi}}}}(\mathbf{x})$. Due to the data-fitting constraint $\mathbf{x} = \mathbf{f}_{\hat{\boldsymbol{\theta}}}(\mathbf{g}_{\hat{\boldsymbol{\phi}}}(\mathbf{x}))$ and the chain rule, we have

$$\mathbf{J}_{\mathbf{f}_{\hat{\boldsymbol{\theta}}}}(\hat{\mathbf{s}})\mathbf{J}_{\mathbf{g}_{\hat{\boldsymbol{\phi}}}}(\mathbf{x}) = \mathbf{I}. \quad (19)$$

Note that $\mathbf{J}_{\mathbf{f}_{\hat{\boldsymbol{\theta}}}}(\hat{\mathbf{s}})$ is full rank, due to the optimality of the problem (7). Therefore, $\mathbf{J}_{\mathbf{g}_{\hat{\boldsymbol{\phi}}}}(\mathbf{x})$ must also be full rank. As a result, $\mathbf{J}_{\mathbf{h}}(\mathbf{s})$ is full rank, and hence $\mathbf{h}$ is indeed a locally invertible function at $\mathbf{s}$, by the inverse function theorem [69, Theorem 6.26].

Now, we can rewrite (17) as

$$(\mathbf{J}_{\mathbf{h}}(\mathbf{s})^{-1})^\top \mathbf{J}_{\mathbf{f}}(\mathbf{s})^\top = \mathbf{J}_{\mathbf{f}_{\hat{\boldsymbol{\theta}}}}(\hat{\mathbf{s}})^\top. \quad (20)$$

By the inverse function theorem [69, Theorem 6.26] and the chain rule, we also have $\mathbf{J}_{\mathbf{h}}(\mathbf{s})^{-1} = \mathbf{J}_{\mathbf{h}^{-1}}(\hat{\mathbf{s}})$, which results in

$$(\mathbf{J}_{\mathbf{h}^{-1}}(\hat{\mathbf{s}}))^\top \mathbf{J}_{\mathbf{f}}(\mathbf{s})^\top = \mathbf{J}_{\mathbf{f}_{\hat{\boldsymbol{\theta}}}}(\hat{\mathbf{s}})^\top. \quad (21)$$

We note that there is an intrinsic scaling ambiguity of the Jacobian matrices here: with an diagonal matrix $\mathbf{D}_{\mathbf{w}}(\mathbf{s}) = \text{diag}(1/w_1(\mathbf{s}), \dots, 1/w_d(\mathbf{s}))$ whose entries are dependent on the unknown $\mathbf{w}(\mathbf{s})$, we have

$$\mathbf{J}_{\mathbf{f}_{\hat{\boldsymbol{\theta}}}}(\hat{\mathbf{s}})^\top = (\mathbf{J}_{\mathbf{h}^{-1}}(\hat{\mathbf{s}}))^\top \mathbf{D}_{\mathbf{w}}(\mathbf{s})^{-1} \mathbf{D}_{\mathbf{w}}(\mathbf{s}) \mathbf{J}_{\mathbf{f}}(\mathbf{s})^\top. \quad (22)$$

We use the shorthand $\mathbf{z}_i := \mathbf{D}_{\mathbf{w}}(\mathbf{s}) \nabla f_i(\mathbf{s})$, which is the i -th column of $\mathbf{D}_{\mathbf{w}}(\mathbf{s}) \mathbf{J}_{\mathbf{f}}(\mathbf{s})^\top$. Notice that $\mathbf{D}_{\mathbf{w}}(\mathbf{s}) \mathbf{J}_{\mathbf{f}}(\mathbf{s})^\top$ is a matrix whose columns lie inside the unit L_1 -norm ball $\mathcal{B}_1$, i.e.

$$\mathbf{z}_1, \dots, \mathbf{z}_m \in \mathcal{B}_1. \quad (23)$$

By combining Proposition A.1 and Assumption 3.1, we have a rescaled version of the SDI condition:

$$\mathcal{B}_2(\frac{1}{\sqrt{d}}) \subset \text{conv}\{\mathbf{z}_1, \dots, \mathbf{z}_m\} \subset \mathcal{B}_1, \text{ and} \quad (24)$$

$$\text{conv}\{\mathbf{z}_1, \dots, \mathbf{z}_m\}^* \cap \text{bd}(\mathcal{B}_2(\sqrt{d})) = \text{extr}(\mathcal{B}_\infty). \quad (25)$$

To proceed, let $\mathbf{H} := (\mathbf{J}_{h^{-1}}(\hat{\mathbf{s}}))^\top \mathbf{D}_w(\mathbf{s})^{-1}$. To show that $\mathbf{J}_h(\mathbf{s})$ has the support structure of a permutation matrix as desired, we can prove that such structure occurs in $\mathbf{H}$. Now, we consider the constraint (7b),

$$\|\mathbf{J}_{f_{\hat{\theta}}}(\hat{\mathbf{s}})_{i,:}\|_1 \leq C, \forall i \in [m]. \quad (26)$$

Combining (26) with the equation (22), we have

$$\|\frac{1}{C}\mathbf{H}\mathbf{z}_i\|_1 \leq 1, \forall i = 1, \dots, m. \quad (27)$$

Observe that (27) can be further rewritten as the following set of 2^d explicit inequalities over all possible signs of each entry of $\frac{1}{C}\mathbf{H}\mathbf{z}_i$:

$$\mathbf{z}_i^\top (\frac{1}{C}\mathbf{H}^\top \mathbf{u}) \leq 1, \forall \mathbf{u} \in \{\pm 1\}^d. \quad (28)$$

Note that $\{\pm 1\}^d = \text{extr}(\mathcal{B}_\infty)$. This allows us to reveal from (28) that

$$\mathbf{z}_i^\top (\frac{1}{C}\mathbf{H}^\top \mathbf{u}) \leq 1, \forall \mathbf{u} \in \text{extr}(\mathcal{B}_\infty). \quad (29)$$

Hence, by definition of the polar set of a polytope, we have

$$\frac{1}{C}\mathbf{H}^\top \mathbf{u} \in \text{conv}\{\mathbf{z}_1, \dots, \mathbf{z}_m\}^*, \forall \mathbf{u} \in \text{extr}(\mathcal{B}_\infty). \quad (30)$$

To proceed, recall from the rescaled SDI assumption in (24) that $\mathcal{B}_2(\frac{1}{\sqrt{d}}) \subset \text{conv}\{\mathbf{z}_1, \dots, \mathbf{z}_m\}$. By Proposition A.3 of polar sets,

$$\text{conv}\{\mathbf{z}_1, \dots, \mathbf{z}_m\}^* \subset \mathcal{B}_2(\sqrt{d}). \quad (31)$$

Since $\frac{1}{C}\mathbf{H}^\top \mathbf{u} \in \text{conv}\{\mathbf{z}_1, \dots, \mathbf{z}_m\}^*$ and $\text{conv}\{\mathbf{z}_1, \dots, \mathbf{z}_m\}^* \subset \mathcal{B}_2(\sqrt{d})$, we have $\frac{1}{C}\mathbf{H}^\top \mathbf{u} \in \mathcal{B}_2(\sqrt{d})$, i.e. $\|\frac{1}{C}\mathbf{H}^\top \mathbf{u}\|_2 \leq \sqrt{d}$ for any $\mathbf{u} \in \text{extr}(\mathcal{B}_\infty)$. By Lemma B.1, we can see that

$$|\det \mathbf{H}| \leq C, \quad (32)$$

with equality holding if and only if $\frac{1}{C}\mathbf{H}$ is an orthogonal matrix. Here, we note that the J-VolMax criterion would lead to an $\mathbf{H}$ with maximal determinant, so then $|\det \mathbf{H}| = C$.

Since $\frac{1}{C}\mathbf{H}$ is an orthogonal matrix and $\|\mathbf{u}\|_2 = \sqrt{d}$ for any $\mathbf{u} \in \text{extr}(\mathcal{B}_\infty)$, we have $\frac{1}{C}\mathbf{H}^\top \mathbf{u} \in \text{bd}(\mathcal{B}_2(\sqrt{d}))$. Also, recall $\frac{1}{C}\mathbf{H}^\top \mathbf{u} \in \text{conv}\{\mathbf{z}_1, \dots, \mathbf{z}_m\}^*$. Hence, $\frac{1}{C}\mathbf{H}^\top \mathbf{u} \in \text{conv}\{\mathbf{z}_1, \dots, \mathbf{z}_m\}^* \cap \text{bd}(\mathcal{B}_2(\sqrt{d}))$. Since

$$\text{conv}\{\mathbf{z}_1, \dots, \mathbf{z}_m\}^* \cap \text{bd}(\mathcal{B}_2(\sqrt{d})) = \text{extr}(\mathcal{B}_\infty), \quad (33)$$

we have

$$\frac{1}{C}\mathbf{H}^\top \mathbf{u} \in \text{extr}(\mathcal{B}_\infty), \forall \mathbf{u} \in \text{extr}(\mathcal{B}_\infty), \quad (34)$$

which implies

$$\|\frac{1}{C}\mathbf{H}_{:,i}\|_1 = 1, \forall i = 1, \dots, d. \quad (35)$$

Hence, $|\det \mathbf{H}| = C$ if and only if both $\|\frac{1}{C}\mathbf{H}_{:,i}\|_1 = 1$ and $\frac{1}{C}\mathbf{H}$ is an orthogonal matrix. Then, we are able to conclude that

$$|\det \mathbf{H}| = C \quad (36)$$

if and only if $\frac{1}{C}\mathbf{H}$ is a signed permutation matrix. Equivalently,

$$\log |\det \mathbf{H}| \leq \log C \quad (37)$$

with equality holds if and only if $\frac{1}{C}\mathbf{H}$ is a signed permutation matrix. The arguments in Eqs. (26)-(37) are from analytical tools first developed in the PMF work [30]. Similar steps were used in [31, 33] for other models as well.

Suppose that there exists an optimal solution $(\bar{\theta}, \bar{\phi})$ with some estimated source $\bar{s} = \mathbf{g}_{\bar{\phi}}(x)$ such that the mapping $\bar{\mathbf{h}} = \mathbf{g}_{\bar{\phi}} \circ \mathbf{f}$ is not a composition of a permutation and an element-wise invertible transformation with strictly positive probability, i.e., for $\bar{\mathbf{H}} = (\mathbf{J}_{\bar{\mathbf{h}}^{-1}}(\bar{s}))^\top \mathbf{D}_w(s)^{-1}$, we have

$$\mathbb{P}(\log |\det \bar{\mathbf{H}}| < \log C) > 0, \quad (38)$$

where the probability is over p_s . This implies

$$\mathbb{E}[\log \det(\mathbf{J}_{\mathbf{f}_{\bar{\theta}}}(\bar{s}))^\top \mathbf{J}_{\mathbf{f}_{\bar{\theta}}}(\bar{s})] \quad (39)$$

$$= \mathbb{E}[\log \det(\mathbf{J}_{\bar{\mathbf{h}}^{-1}}(\bar{s}))^\top \mathbf{J}_{\mathbf{f}}(s)^\top \mathbf{J}_{\mathbf{f}}(s) \mathbf{J}_{\bar{\mathbf{h}}^{-1}}(\bar{s})] \quad (40)$$

$$= \mathbb{E}[\log \det(\mathbf{J}_{\bar{\mathbf{h}}^{-1}}(\bar{s}))^\top \mathbf{D}_w(s)^{-1} \mathbf{D}_w(s) \mathbf{J}_{\mathbf{f}}(s)^\top \mathbf{J}_{\mathbf{f}}(s) \mathbf{D}_w(s) \mathbf{D}_w(s)^{-1} \mathbf{J}_{\bar{\mathbf{h}}^{-1}}(\bar{s})] \quad (41)$$

$$= \mathbb{E}[\log \det(\bar{\mathbf{H}} \mathbf{D}_w(s) \mathbf{J}_{\mathbf{f}}(s)^\top \mathbf{J}_{\mathbf{f}}(s) \mathbf{D}_w(s) \bar{\mathbf{H}}^\top)] \quad (42)$$

$$= \mathbb{E}[\log(\det(\bar{\mathbf{H}})^2 \det(\mathbf{D}_w(s) \mathbf{J}_{\mathbf{f}}(s)^\top \mathbf{J}_{\mathbf{f}}(s) \mathbf{D}_w(s)))] \quad (43)$$

$$= \mathbb{E}[2 \log |\det \bar{\mathbf{H}}| + \log \det(\mathbf{D}_w(s) \mathbf{J}_{\mathbf{f}}(s)^\top \mathbf{J}_{\mathbf{f}}(s) \mathbf{D}_w(s))] \quad (44)$$

$$\stackrel{(a)}{=} \int_{s \in \mathcal{A}} p(s) [2 \log |\det \bar{\mathbf{H}}| + \log \det(\mathbf{D}_w(s) \mathbf{J}_{\mathbf{f}}(s)^\top \mathbf{J}_{\mathbf{f}}(s) \mathbf{D}_w(s))] ds \quad (45)$$

$$+ \int_{s \in \mathcal{S} \setminus \mathcal{A}} p(s) [2 \log |\det \bar{\mathbf{H}}| + \log \det(\mathbf{D}_w(s) \mathbf{J}_{\mathbf{f}}(s)^\top \mathbf{J}_{\mathbf{f}}(s) \mathbf{D}_w(s))] ds \quad (46)$$

$$\stackrel{(b)}{<} \int_{s \in \mathcal{A}} p(s) [2 \log C + \log \det(\mathbf{D}_w(s) \mathbf{J}_{\mathbf{f}}(s)^\top \mathbf{J}_{\mathbf{f}}(s) \mathbf{D}_w(s))] ds \quad (47)$$

$$+ \int_{s \in \mathcal{S} \setminus \mathcal{A}} p(s) [2 \log C + \log \det(\mathbf{D}_w(s) \mathbf{J}_{\mathbf{f}}(s)^\top \mathbf{J}_{\mathbf{f}}(s) \mathbf{D}_w(s))] ds \quad (48)$$

$$= \int_{s \in \mathcal{S}} p(s) [2 \log C + \log \det(\mathbf{D}_w(s) \mathbf{J}_{\mathbf{f}}(s)^\top \mathbf{J}_{\mathbf{f}}(s) \mathbf{D}_w(s))] ds \quad (49)$$

$$= \mathbb{E}[2 \log C + \log \det(\mathbf{D}_w(s) \mathbf{J}_{\mathbf{f}}(s)^\top \mathbf{J}_{\mathbf{f}}(s) \mathbf{D}_w(s))], \quad (50)$$

$$= \mathbb{E}[\log(C^2 \det(\mathbf{D}_w(s) \mathbf{J}_{\mathbf{f}}(s)^\top \mathbf{J}_{\mathbf{f}}(s) \mathbf{D}_w(s)))] \quad (51)$$

where (a) and (b) involve two sets $\mathcal{A} = \{s : \log |\det \bar{\mathbf{H}}| < C\}$ and $\mathcal{S} \setminus \mathcal{A} = \{s : \log |\det \bar{\mathbf{H}}| = C\}$.

Consider an invertible continuous element-wise function $\tilde{\mathbf{h}} : \mathcal{S} \mapsto \mathcal{S}$ with $\tilde{s} := \tilde{\mathbf{h}}^{-1}(s) \in \mathcal{S}$, $\tilde{\mathbf{h}}(\tilde{s}) = s$, and $\mathbf{J}_{\tilde{\mathbf{h}}}(\tilde{s}) = \mathbf{J}_{\tilde{\mathbf{h}}}(\tilde{\mathbf{h}}^{-1}(s)) = C \mathbf{D}_w(s)$. This function $\tilde{\mathbf{h}}$ merely rescales and maps each latent component by an element-wise invertible transformation. Define

$$\tilde{\mathbf{f}}(\tilde{s}) := \mathbf{f}(\tilde{\mathbf{h}}(\tilde{s})), \quad (52)$$

whose Jacobian is

$$\mathbf{J}_{\tilde{\mathbf{f}}}(\tilde{s}) = \mathbf{J}_{\mathbf{f}}(s) \mathbf{J}_{\tilde{\mathbf{h}}}(\tilde{s}). \quad (53)$$

We show that $\tilde{\mathbf{f}}(\tilde{s})$ is a feasible solution to J-VolMax, i.e. satisfying (7b) and (7c). First, note that

$$\|\mathbf{J}_{\tilde{\mathbf{f}}}(\tilde{s})_{i,:}\|_1 = \|C \mathbf{D}_w(s) \nabla f_i(s)\|_1 \leq C. \quad (54)$$

Second, observe that

$$\tilde{\mathbf{f}}(\tilde{s}) = \mathbf{f}(\tilde{\mathbf{h}}(\tilde{s})) = \mathbf{f}(\tilde{\mathbf{h}}(\tilde{\mathbf{h}}^{-1}(s))) = \mathbf{f}(s) = x. \quad (55)$$

Hence, there exists a feasible solution $x = \tilde{\mathbf{f}}(\tilde{s})$ to J-VolMax (7) such that

$$\mathbb{E}[\log \det(\mathbf{J}_{\mathbf{f}_{\bar{\theta}}}(\bar{s}))^\top \mathbf{J}_{\mathbf{f}_{\bar{\theta}}}(\bar{s})] < \mathbb{E}[\log \det(\mathbf{J}_{\tilde{\mathbf{f}}}(\tilde{s}))^\top \mathbf{J}_{\tilde{\mathbf{f}}}(\tilde{s})]. \quad (56)$$

This contradicts the optimality of $(\bar{\theta}, \bar{\phi})$ to J-VolMax problem in (7). We hence conclude that any optimal solution (θ^*, ϕ^*) of (7) would have the composite function $\mathbf{h}^* = \mathbf{g}_{\phi^*} \circ \mathbf{f}$ satisfying

$$\mathbf{J}_{\mathbf{h}^*}(s) = \mathbf{D}(s) \mathbf{\Pi}(s) \text{ a.e.,} \quad (57)$$

where $\mathbf{D}(s)$ is an invertible diagonal matrix such that with the i^{th} diagonal entries dependent on s_i due to the definition of a Jacobian matrix, and $\mathbf{\Pi}(s)$ is a permutation matrix that depends on s . $\square$

Theorem 3.2 (Identifiability of J-VolMax). *Denote any optimal solution of Problem (7) as $(\hat{\theta}, \hat{\phi})$. Assume $\hat{f} = f_{\hat{\theta}}$ and $\hat{g} = g_{\hat{\phi}}$ are universal function representers. Suppose the model in (1) and Assumption 3.1 hold for every $s \in \mathcal{S}$. Then, we have $\hat{s} = \hat{g}(x) = \hat{g} \circ f(s)$ where*

$$[\hat{s}]_i = [\hat{g}(x)]_i = \rho_i(s_{\pi(i)}), \forall i \in [d], \quad (8)$$

in which π is a permutation of $\{1, \dots, d\}$ and $\rho_i(\cdot) : \mathbb{R} \rightarrow \mathbb{R}$ is an invertible function.

Proof of Theorem 3.2. Define a differentiable function $h^* : \mathcal{S} \mapsto \mathcal{S}$ as a composition $h = \hat{g} \circ f$. By Lemma B.2, we have that the Jacobian at the point $s \in \mathcal{S}$ of the function $h^* = \hat{g} \circ f$ will satisfy

$$J_{h^*}(s) = D(s)\Pi(s), \quad (58)$$

where $D(s)$ is an invertible diagonal matrix with $[D(s)]_{i,i}$ dependent on i -th component s_i only, and $\Pi(s)$ is a permutation matrix dependent on the point s , *almost everywhere*.

First, we want to leverage continuity of h^* to show that $J_{h^*}(s) = D(s)\Pi(s)$ actually holds everywhere on $\mathcal{S}$. This is because if there exists $\bar{s} \in \mathcal{S}$ such that $J_{h^*}(\bar{s}) \neq D(\bar{s})\Pi(\bar{s})$, then due to the full-rank property of $J_{h^*}(\bar{s})$, there exist a $j \in [d]$ such that the column $[J_{h^*}(\bar{s})]_{:j}$ has at least two non-zero elements. However, the continuity of the Jacobian J_{h^*} implies that there exists an open neighborhood of $\bar{s}$, where the column $[J_{h^*}(\bar{s})]_{:j}$ has at least two non-zero elements. This contradicts the fact that $J_{h^*}(s) = D(s)\Pi(s)$ almost everywhere.

Next, it remains to show that $\Pi(s)$ is constant (say $\Pi \in \mathcal{P}_d$) for all $s \in \mathcal{S}$. The reason is that if the non-zero elements in $\Pi(s)$ switched locations, then there would exist a point in $\bar{s} \in \mathcal{S}$ where $J_{h^*}(\bar{s})$ is singular, which contradicts the full-rank property of $J_{h^*}(\bar{s})$. Hence, $\Pi(s) = \Pi, \forall s \in \mathcal{S}$; see a similar argument in [12, Theorem 1].

Finally, let π denote the permutation of $\{1, \dots, d\}$ corresponding to the permutation matrix Π , i.e., let $r = [1, 2, \dots, d]^\top$, then $\pi(i) = \Pi_{i:r}$. Then, $J_{h^*}(s) = D(s)\Pi$ implies that

$$\frac{\partial h_i^*(s)}{\partial s_{\pi(i)}} = \frac{\partial \hat{s}_i}{\partial s_{\pi(i)}} \neq 0 \quad \text{and} \quad \frac{\partial \hat{s}_i}{\partial s_{\pi(j)}} = 0, \forall j \neq i$$

Hence, there exist scalar functions $\rho_1, \dots, \rho_d$ such that

$$\hat{s}_i = \rho_i(s_{\pi(i)}).$$

Note that $\rho_1, \dots, \rho_d$ are invertible functions by the invertibility of h^* , which was shown in Lemma B.2. In conclusion, the estimated unmixer g_{ϕ^*} satisfies $g_{\phi^*} \circ f$ being an invertible element-wise transformation and permutation. $\square$

C Proof of Theorem 3.3

Before diving into proving Theorem 3.3, we first show three cornerstone lemmas of the proof. These lemmas help characterize the detailed conditions under which an optimal solution of J-VolMax criterion can recover the latent components up to a constant permutation and invertible element-wise transformations, i.e.,

$$\hat{g}(x^{(n)}) = \hat{\Pi}\hat{\rho}(s^{(n)}), \forall s^{(n)} \in \mathcal{S}_N. \quad (59)$$

The formal result is given in Lemma C.3. Intuitively, (59) holds if the points $s^{(n)}$ in set $\mathcal{S}_N$ are sampled densely and closely enough from p_s with a sufficiently large N , together with some regularity conditions on the ground-truth and the learnable function classes. To show (59), we employ a three-step argument, Lemma C.1, Lemma C.2, and Lemma C.3.

Firstly, we show that for any feasible solution of J-VolMax that attains maximal Jacobian volume at each $s^{(n)} \in \mathcal{S}_N$, it must give an estimated latent component vector that is permutation and invertible element-wise transformation of the ground truth $s^{(n)}$. In fact, due to continuity, this holds at every point s in a union of neighborhoods $U_N = \cup_{n=1}^N U^{(n)}$ centered on each $s^{(n)}$. The result is stated in the following Lemma C.1, and proof is a straightforward adaptation of Lemma B.2.

Lemma C.1. Denote any feasible solution of J-VolMax problem (7) as $\bar{\Theta} := [\bar{\theta}, \bar{\phi}]$, learned from classes of learnable encoders and decoders $\mathcal{F}, \mathcal{G}$ that include the function classes of ground-truth encoders and decoders $\mathcal{F}', \mathcal{G}'$. Assume that there is an unknown finite set $\mathcal{S}_N := \{\mathbf{s}^{(1)}, \dots, \mathbf{s}^{(N)}\} \subset \mathcal{S}$ with unknown $\mathcal{X}_N := \{\mathbf{x}^{(n)} \in \mathcal{X} : \mathbf{x}^{(n)} = \mathbf{f}(\mathbf{s}^{(n)}), \forall \mathbf{s}^{(n)} \in \mathcal{S}_N\} \subset \mathcal{X}$ such that the Assumption 3.1 is satisfied at each of the N points in $\mathcal{S}_N$. If the estimated latent components $\bar{\mathbf{s}}^{(n)} = \mathbf{g}_{\bar{\phi}}(\mathbf{x}^{(n)})$ are not permutation and invertible element-wise transformation of ground-truth $\mathbf{s}^{(n)}$, then there exists another feasible solution with parameters $\tilde{\Theta}$ such that

$$V(\mathbf{s}^{(n)}; \bar{\Theta}) < V(\mathbf{s}^{(n)}; \tilde{\Theta}), \forall \mathbf{s}^{(n)} \in \mathcal{S}_N \quad (60)$$

Furthermore, there exists an union $U_N = \cup_{n=1}^N U^{(n)} \subseteq \mathcal{S}$ of open ball neighborhoods $U^{(n)} = \{\mathbf{s} \in \mathcal{S} : \|\mathbf{s} - \mathbf{s}^{(n)}\|_2 < d^{(n)}\}$ centered on $\mathbf{s}^{(n)} \in \mathcal{S}_N$ such that

$$V(\mathbf{s}; \bar{\Theta}) < V(\mathbf{s}; \tilde{\Theta}), \forall \mathbf{s} \in U_N. \quad (61)$$

Proof. For the feasible solution $\bar{\mathbf{f}} = \mathbf{f}_{\bar{\theta}}$ and $\bar{\mathbf{g}} = \mathbf{g}_{\bar{\phi}}$, define a differentiable function $\bar{\mathbf{h}} : \mathcal{S} \mapsto \mathcal{S}$ as a composition $\bar{\mathbf{h}} = \bar{\mathbf{g}} \circ \mathbf{f}$, and consider a data point $\mathbf{x}^{(n)} = \mathbf{f}(\mathbf{s}^{(n)}) \in \mathcal{X}_N$ for $\mathbf{s}^{(n)} \in \mathcal{S}_N$. Let $\bar{\mathbf{s}}^{(n)} = \bar{\mathbf{g}}(\mathbf{x}^{(n)})$. Following the same argument (16)-(36) as in Lemma B.2, for $\bar{\mathbf{H}} := (\mathbf{J}_{\bar{\mathbf{h}}^{-1}}(\bar{\mathbf{s}}^{(n)})^\top \mathbf{D}_{\mathbf{w}}(\mathbf{s}^{(n)})^{-1})$, we can similarly derive that $\log |\det \bar{\mathbf{H}}| \leq \log C$ at any point $\mathbf{s}^{(n)} \in \mathcal{S}_N$, with equality holds if and only if $\frac{1}{C} \bar{\mathbf{H}}$ is a signed permutation matrix, i.e. when $\bar{\mathbf{h}}$ is a composition of a permutation and invertible element-wise transformations.

We will show in the following that if $\bar{\mathbf{h}}$ is not a composition of a permutation and an element-wise invertible transformation, then the Jacobian volume $V(\mathbf{s}^{(n)}; \bar{\Theta})$ attained at $\mathbf{s}^{(n)}$ of the feasible solution $\bar{\Theta}$ is strictly smaller than the Jacobian volume of another feasible solution; hence, the feasible solution $\bar{\Theta}$ does not reach maximal Jacobian volume at $\mathbf{s}^{(n)}$.

Since $\bar{\mathbf{h}} = \bar{\mathbf{g}} \circ \mathbf{f}$ is not a composition of permutation and component-wise invertible mappings,

$$\log |\det \bar{\mathbf{H}}| < \log C, \forall \mathbf{s}^{(n)} \in \mathcal{S}_N. \quad (62)$$

This implies

$$\begin{aligned} & \log \det(\mathbf{J}_{\bar{\mathbf{f}}_{\bar{\theta}}}(\bar{\mathbf{s}}^{(n)})^\top \mathbf{J}_{\bar{\mathbf{f}}_{\bar{\theta}}}(\bar{\mathbf{s}}^{(n)})) \\ &= \log \det(\mathbf{J}_{\bar{\mathbf{h}}^{-1}}(\bar{\mathbf{s}}^{(n)})^\top \mathbf{J}_{\mathbf{f}}(\mathbf{s}^{(n)})^\top \mathbf{J}_{\mathbf{f}}(\mathbf{s}^{(n)}) \mathbf{J}_{\bar{\mathbf{h}}^{-1}}(\bar{\mathbf{s}}^{(n)})) \\ &= \log \det(\mathbf{J}_{\bar{\mathbf{h}}^{-1}}(\bar{\mathbf{s}}^{(n)})^\top \mathbf{D}_{\mathbf{w}}(\mathbf{s}^{(n)})^{-1} \mathbf{D}_{\mathbf{w}}(\mathbf{s}^{(n)}) \mathbf{J}_{\mathbf{f}}(\mathbf{s}^{(n)})^\top \mathbf{J}_{\mathbf{f}}(\mathbf{s}^{(n)}) \mathbf{D}_{\mathbf{w}}(\mathbf{s}^{(n)}) \mathbf{D}_{\mathbf{w}}(\mathbf{s}^{(n)})^{-1} \mathbf{J}_{\bar{\mathbf{h}}^{-1}}(\bar{\mathbf{s}}^{(n)})) \\ &= \log \det(\bar{\mathbf{H}} \mathbf{D}_{\mathbf{w}}(\mathbf{s}^{(n)}) \mathbf{J}_{\mathbf{f}}(\mathbf{s}^{(n)})^\top \mathbf{J}_{\mathbf{f}}(\mathbf{s}^{(n)}) \mathbf{D}_{\mathbf{w}}(\mathbf{s}^{(n)}) \bar{\mathbf{H}}^\top) \\ &= \log[(\det \bar{\mathbf{H}})^2 \det(\mathbf{D}_{\mathbf{w}}(\mathbf{s}^{(n)}) \mathbf{J}_{\mathbf{f}}(\mathbf{s}^{(n)})^\top \mathbf{J}_{\mathbf{f}}(\mathbf{s}^{(n)}) \mathbf{D}_{\mathbf{w}}(\mathbf{s}^{(n)}))] \\ &< \log(C^2 \det(\mathbf{D}_{\mathbf{w}}(\mathbf{s}^{(n)}) \mathbf{J}_{\mathbf{f}}(\mathbf{s}^{(n)})^\top \mathbf{J}_{\mathbf{f}}(\mathbf{s}^{(n)}) \mathbf{D}_{\mathbf{w}}(\mathbf{s}^{(n)}))), \forall \mathbf{s}^{(n)} \in \mathcal{S}_N. \end{aligned} \quad (63)$$

Consider an invertible continuous element-wise function $\tilde{\mathbf{h}} : \mathcal{S} \mapsto \mathcal{S}$ with $\tilde{\mathbf{s}}^{(n)} := \tilde{\mathbf{h}}^{-1}(\mathbf{s}^{(n)}) \in \mathcal{S}$, $\tilde{\mathbf{h}}(\tilde{\mathbf{s}}^{(n)}) = \mathbf{s}^{(n)}$, and $\mathbf{J}_{\tilde{\mathbf{h}}}(\tilde{\mathbf{s}}^{(n)}) = \mathbf{J}_{\tilde{\mathbf{h}}}(\tilde{\mathbf{h}}^{-1}(\mathbf{s}^{(n)})) = C \mathbf{D}_{\mathbf{w}}(\mathbf{s}^{(n)})$. This function $\tilde{\mathbf{h}}$ merely rescales and maps each latent component by an element-wise invertible transformation. Define

$$\tilde{\mathbf{f}}(\tilde{\mathbf{s}}^{(n)}) := \mathbf{f}(\tilde{\mathbf{h}}(\tilde{\mathbf{s}}^{(n)})), \quad (64)$$

whose Jacobian is

$$\mathbf{J}_{\tilde{\mathbf{f}}}(\tilde{\mathbf{s}}^{(n)}) = \mathbf{J}_{\mathbf{f}}(\mathbf{s}^{(n)}) \mathbf{J}_{\tilde{\mathbf{h}}}(\tilde{\mathbf{s}}^{(n)}). \quad (65)$$

We show that $\tilde{\mathbf{f}}(\tilde{\mathbf{s}}^{(n)})$ is a feasible solution to J-VolMax (7), i.e. (7b) and (7c) hold. First, note that

$$\|\mathbf{J}_{\tilde{\mathbf{f}}}(\tilde{\mathbf{s}}^{(n)})_{i,:}\|_1 = \|C \mathbf{D}_{\mathbf{w}}(\mathbf{s}^{(n)}) \nabla f_i(\mathbf{s}^{(n)})\|_1 \leq C, \forall i \in [m]. \quad (66)$$

Second, observe that

$$\tilde{\mathbf{f}}(\tilde{\mathbf{s}}^{(n)}) = \mathbf{f}(\tilde{\mathbf{h}}(\tilde{\mathbf{s}}^{(n)})) = \mathbf{f}(\tilde{\mathbf{h}}(\tilde{\mathbf{h}}^{-1}(\mathbf{s}^{(n)}))) = \mathbf{f}(\mathbf{s}^{(n)}) = \mathbf{x}^{(n)}. \quad (67)$$

Hence, there exists a feasible solution $\mathbf{x} = \tilde{\mathbf{f}}(\tilde{\mathbf{s}}^{(n)})$ to J-VolMax such that

$$\log \det(\mathbf{J}_{\tilde{\mathbf{f}}_{\tilde{\boldsymbol{\theta}}}}(\tilde{\mathbf{s}}^{(n)})^\top \mathbf{J}_{\tilde{\mathbf{f}}_{\tilde{\boldsymbol{\theta}}}}(\tilde{\mathbf{s}}^{(n)})) < \log \det(\mathbf{J}_{\tilde{\mathbf{f}}}(\tilde{\mathbf{s}}^{(n)})^\top \mathbf{J}_{\tilde{\mathbf{f}}}(\tilde{\mathbf{s}}^{(n)})), \forall \mathbf{s}^{(n)} \in \mathcal{S}_N. \quad (68)$$

Equivalently, for $\tilde{\Theta} = [\tilde{\boldsymbol{\theta}}, \tilde{\boldsymbol{\phi}}]$ being the parameters corresponding to $\tilde{\mathbf{f}}$ and $\tilde{\mathbf{s}}^{(n)}$,

$$V(\mathbf{s}^{(n)}; \bar{\Theta}) < V(\mathbf{s}^{(n)}; \tilde{\Theta}), \forall \mathbf{s}^{(n)} \in \mathcal{S}_N. \quad (69)$$

Finally, due to continuity of the involved functions, there exists a neighborhood $U^{(n)} = \{\mathbf{s} \in \mathcal{S} : \|\mathbf{s} - \mathbf{s}^{(n)}\|_2 < d^{(n)}\}$ centered at $\mathbf{s}^{(n)}$ such that

$$V(\mathbf{s}; \bar{\Theta}) < V(\mathbf{s}; \tilde{\Theta}), \forall \mathbf{s} \in U^{(n)}, \quad (70)$$

which holds for all $U^{(1)}, \dots, U^{(n)}$. Hence, with $U_N = \cup_{n=1}^N U^{(n)}$,

$$V(\mathbf{s}; \bar{\Theta}) < V(\mathbf{s}; \tilde{\Theta}), \forall \mathbf{s} \in U_N. \quad (71)$$

□

Following Lemma C.1, we show that if the SDI-satisfying finite set $\mathcal{S}_N$ are sampled dense enough such that U_N covers a significant part of $\mathcal{S}$, then any optimal solution of J-VolMax criterion must recover the ground-truth latent components $\mathbf{s}^{(n)}$, up to $\mathbf{s}^{(n)}$ -dependent permutation and invertible element-wise transformations. The proof leverages the result from the previous Lemma C.1 to show a contradiction that if a feasible solution $\bar{\Theta}$ does not give us ground-truth latents up to aforementioned ambiguities, there must exist another feasible solution $\tilde{\Theta}$ that is strictly better than $\bar{\Theta}$ in terms of J-VolMax objective (7a), i.e.,

$$\mathbb{E}[V(\mathbf{s}; \tilde{\Theta})] > \mathbb{E}[V(\mathbf{s}; \bar{\Theta})]. \quad (72)$$

The result is formally given in the following Lemma C.2.

Lemma C.2. Denote any optimal solution of J-VolMax problem (7) as $\hat{\Theta} := [\hat{\boldsymbol{\theta}}, \hat{\boldsymbol{\phi}}]$, learned from classes of learnable encoders and decoders $\mathcal{F}, \mathcal{G}$ that include the function classes of ground-truth encoders and decoders $\mathcal{F}', \mathcal{G}'$. Suppose there is an unknown finite set $\mathcal{S}_N := \{\mathbf{s}^{(1)}, \dots, \mathbf{s}^{(N)}\} \subset \mathcal{S}$ with unknown $\mathcal{X}_N := \{\mathbf{x}^{(n)} \in \mathcal{X} : \mathbf{x}^{(n)} = \mathbf{f}(\mathbf{s}^{(n)}), \forall \mathbf{s}^{(n)} \in \mathcal{S}_N\} \subset \mathcal{X}$ such that the Assumption 3.1 is satisfied at each of the N points in $\mathcal{S}_N$.

For the union of neighborhoods $U_N = \cup_{n=1}^N U^{(n)} = \cup_{n=1}^N \{\mathbf{s} \in \mathcal{S} : \|\mathbf{s} - \mathbf{s}^{(n)}\|_2 < d^{(n)}\}$ within which $V(\mathbf{s}; \hat{\Theta})$ is optimal, assume that $\mathcal{S}_N$ is sampled densely enough from $p_{\mathbf{s}}$ with sufficiently large N such that

$$\frac{\mathbb{P}(\mathbf{s} \in U_N)}{\mathbb{P}(\mathbf{s} \in \mathcal{S} \setminus U_N)} > \frac{G_{\max}}{G_{\min}} > 1, \quad (73)$$

where $G_{\min}, G_{\max}$ are bi-Lipschitz constants of the Jacobian volume surrogate with respect to the parameter, i.e.,

$$G_{\min} \|\Theta_1 - \Theta_2\|_2 \leq |V(\mathbf{s}; \Theta_1) - V(\mathbf{s}; \Theta_2)| \leq G_{\max} \|\Theta_1 - \Theta_2\|_2, \forall \mathbf{s} \in \mathcal{S}, \quad (74)$$

for any parameters Θ_1, Θ_2 in $\mathcal{F}, \mathcal{G}$. Then, the optimal encoder $\hat{\mathbf{g}} = \mathbf{g}_{\hat{\boldsymbol{\phi}}}$ gives

$$\hat{\mathbf{g}}(\mathbf{x}^{(n)}) = \hat{\boldsymbol{\Pi}}(\mathbf{s}^{(n)}) \hat{\boldsymbol{\rho}}(\mathbf{s}^{(n)}), \forall \mathbf{s}^{(n)} \in \mathcal{S}_N, \quad (75)$$

where $\hat{\boldsymbol{\Pi}}(\mathbf{s}^{(n)})$ is a $\mathbf{s}^{(n)}$ -dependent permutation matrix, and $\hat{\boldsymbol{\rho}}(\mathbf{s}^{(n)}) = [\hat{\rho}_1(s_1), \dots, \hat{\rho}_d(s_d)]^\top$ is a vector of d invertible element-wise transformations on each component of $\mathbf{s}^{(n)}$.

Proof. Suppose that there exists an optimal solution $\bar{\Theta} = [\bar{\boldsymbol{\theta}}, \bar{\boldsymbol{\phi}}]$ of J-VolMax such that the mapping $\bar{\mathbf{h}} = \mathbf{g}_{\bar{\boldsymbol{\phi}}} \circ \mathbf{f}$ is not a composition of permutation and element-wise invertible transformations. By Lemma C.1, within a certain union of neighborhoods U_N such that (73) is satisfied, there exists another feasible solution $\tilde{\Theta} = [\tilde{\boldsymbol{\theta}}, \tilde{\boldsymbol{\phi}}]$ such that,

$$V(\mathbf{s}; \bar{\Theta}) < V(\mathbf{s}; \tilde{\Theta}), \forall \mathbf{s} \in U_N. \quad (76)$$

We now show that: since the solution $\bar{\Theta}$ does not reach maximal Jacobian volume at each point $\mathbf{s}$ within U_N as in (76), $\bar{\Theta}$ cannot reach maximal expected Jacobian volume over all of $\mathcal{S}$, i.e.,

$$\mathbb{E}[V(\mathbf{s}; \bar{\Theta})] < \mathbb{E}[V(\mathbf{s}; \tilde{\Theta})], \quad (77)$$

and therefore $\bar{\Theta}$ is actually not an optimal solution of J-VolMax, which raises a contradiction.

To begin, notice that the open neighborhood $U^{(n)} = \{\mathbf{s} \in \mathcal{S} : \|\mathbf{s} - \mathbf{s}^{(n)}\| < d^{(n)}\}$ has non-zero measure, then their union $U_N = \cup_{n=1}^N U^{(n)}$ also has non-zero measure. This allows us to define

$$C_{U_N}(\bar{\Theta}) := \int_{\mathbf{s} \in U_N} V(\mathbf{s}; \bar{\Theta}) p(\mathbf{s}) d\mathbf{s}, \quad (78)$$

$$C_{\mathcal{S} \setminus U_N}(\bar{\Theta}) := \int_{\mathbf{s} \in \mathcal{S} \setminus U_N} V(\mathbf{s}; \bar{\Theta}) p(\mathbf{s}) d\mathbf{s}, \quad (79)$$

as the part over U_N and the part over $\mathcal{S} \setminus U_N$ of the objective value $\mathbb{E}[V(\mathbf{s}; \bar{\Theta})]$, i.e., $C_{U_N}(\bar{\Theta}) + C_{\mathcal{S} \setminus U_N}(\bar{\Theta}) = \mathbb{E}[V(\mathbf{s}; \bar{\Theta})]$. Similarly, define

$$C_{U_N}(\tilde{\Theta}) := \int_{\mathbf{s} \in U_N} V(\mathbf{s}; \tilde{\Theta}) p(\mathbf{s}) d\mathbf{s}, \quad (80)$$

$$C_{\mathcal{S} \setminus U_N}(\tilde{\Theta}) := \int_{\mathbf{s} \in \mathcal{S} \setminus U_N} V(\mathbf{s}; \tilde{\Theta}) p(\mathbf{s}) d\mathbf{s}, \quad (81)$$

with $C_{U_N}(\tilde{\Theta}) + C_{\mathcal{S} \setminus U_N}(\tilde{\Theta}) = \mathbb{E}[V(\mathbf{s}; \tilde{\Theta})]$.

Now, observe that

$$C_{U_N}(\tilde{\Theta}) - C_{U_N}(\bar{\Theta}) = \int_{\mathbf{s} \in U_N} (V(\mathbf{s}; \tilde{\Theta}) - V(\mathbf{s}; \bar{\Theta})) p(\mathbf{s}) d\mathbf{s} \quad (82)$$

$$= \int_{\mathbf{s} \in U_N} |V(\mathbf{s}; \tilde{\Theta}) - V(\mathbf{s}; \bar{\Theta})| p(\mathbf{s}) d\mathbf{s} \quad (83)$$

$$\geq \int_{\mathbf{s} \in U_N} G_{\min} \|\tilde{\Theta} - \bar{\Theta}\|_2 p(\mathbf{s}) d\mathbf{s} \quad (84)$$

$$= G_{\min} \|\tilde{\Theta} - \bar{\Theta}\|_2 \int_{\mathbf{s} \in U_N} p(\mathbf{s}) d\mathbf{s} \quad (85)$$

$$= G_{\min} \|\tilde{\Theta} - \bar{\Theta}\|_2 \mathbb{P}(\mathbf{s} \in U_N), \quad (86)$$

where (83) is due to the fact that $V(\mathbf{s}; \tilde{\Theta}) > V(\mathbf{s}; \bar{\Theta})$, $\forall \mathbf{s} \in U_N$ as from Lemma C.1, and (84) is obtained by applying bi-Lipschitz continuity of $V(\mathbf{s}; \cdot)$ as assumed in (74). Similarly, we have

$$|C_{\mathcal{S} \setminus U_N}(\tilde{\Theta}) - C_{\mathcal{S} \setminus U_N}(\bar{\Theta})| \leq \int_{\mathbf{s} \in \mathcal{S} \setminus U_N} |V(\mathbf{s}; \tilde{\Theta}) - V(\mathbf{s}; \bar{\Theta})| p(\mathbf{s}) d\mathbf{s} \quad (87)$$

$$\leq \int_{\mathbf{s} \in \mathcal{S} \setminus U_N} G_{\max} \|\tilde{\Theta} - \bar{\Theta}\|_2 p(\mathbf{s}) d\mathbf{s} \quad (88)$$

$$= G_{\max} \|\tilde{\Theta} - \bar{\Theta}\|_2 \int_{\mathbf{s} \in \mathcal{S} \setminus U_N} p(\mathbf{s}) d\mathbf{s} \quad (89)$$

$$= G_{\max} \|\tilde{\Theta} - \bar{\Theta}\|_2 \mathbb{P}(\mathbf{s} \in \mathcal{S} \setminus U_N), \quad (90)$$

where (87) is obtained via triangle inequality, and (88) is from the bi-Lipschitz continuity of $V(\mathbf{s}; \cdot)$ as in the assumption (74).

Combining (86), (90) with the assumption that $\mathcal{S}_N = \{\mathbf{s}^{(1)}, \dots, \mathbf{s}^{(N)}\}$ is sampled densely and closely enough over $\mathcal{S}$ via $p_{\mathcal{S}}$ such that

$$\frac{\mathbb{P}(\mathbf{s} \in U_N)}{\mathbb{P}(\mathbf{s} \in \mathcal{S} \setminus U_N)} > \frac{G_{\max}}{G_{\min}} > 1, \quad (91)$$

we have the following chain of inequalities:

$$C_{U_N}(\tilde{\Theta}) - C_{U_N}(\bar{\Theta}) \geq G_{\min} \|\tilde{\Theta} - \bar{\Theta}\|_2 \mathbb{P}(\mathbf{s} \in U_N) \quad (92)$$

$$> G_{\max} \|\tilde{\Theta} - \bar{\Theta}\|_2 \mathbb{P}(\mathbf{s} \in \mathcal{S} \setminus U_N) \quad (93)$$

$$\geq |C_{\mathcal{S} \setminus U_N}(\tilde{\Theta}) - C_{\mathcal{S} \setminus U_N}(\bar{\Theta})| \quad (94)$$

$$\geq C_{\mathcal{S} \setminus U_N}(\bar{\Theta}) - C_{\mathcal{S} \setminus U_N}(\tilde{\Theta}). \quad (95)$$

Therefore, the following strict inequality holds:

$$C_{U_N}(\tilde{\Theta}) + C_{\mathcal{S} \setminus U_N}(\tilde{\Theta}) > C_{U_N}(\bar{\Theta}) + C_{\mathcal{S} \setminus U_N}(\bar{\Theta}), \quad (96)$$

or equivalently,

$$\mathbb{E}[V(\mathbf{s}; \tilde{\Theta})] > \mathbb{E}[V(\mathbf{s}; \bar{\Theta})]. \quad (97)$$

That is to say, $\tilde{\Theta} = [\tilde{\theta}, \tilde{\phi}]$ is indeed a feasible solution of J-VolMax that is strictly better than $\bar{\Theta} = (\bar{\theta}, \bar{\phi})$ in terms of expected Jacobian volume in J-VolMax objective (7a). This contradicts the assumed optimality of $(\bar{\theta}, \bar{\phi})$ to J-VolMax problem (7).

We hence conclude that any optimal solution with encoder $\hat{\mathbf{g}}$, the Jacobian at every point $\mathbf{s}^{(n)} \in \mathcal{S}_N$ of the function $\mathbf{h}^* = \hat{\mathbf{g}} \circ \mathbf{f}$ must satisfy

$$\mathbf{J}_{\mathbf{h}^*}(\mathbf{s}^{(n)}) = \mathbf{D}(\mathbf{s}^{(n)})\mathbf{\Pi}(\mathbf{s}^{(n)}),$$

where $\mathbf{D}(\mathbf{s}^{(n)})$ is an invertible diagonal matrix with $[\mathbf{D}(\mathbf{s}^{(n)})]_{i,i}$ dependent on i -th component $s_i^{(n)}$ only, and $\mathbf{\Pi}(\mathbf{s}^{(n)})$ is a permutation matrix dependent on the point $\mathbf{s}^{(n)}$. $\square$

In the following Lemma C.3, we show that under further regularity conditions, the permutation ordering in the result (75) of Lemma C.2,

$$\hat{\mathbf{g}}(\mathbf{x}^{(n)}) = \hat{\mathbf{\Pi}}(\mathbf{s}^{(n)})\hat{\boldsymbol{\rho}}(\mathbf{s}^{(n)}), \quad \forall \mathbf{s}^{(n)} \in \mathcal{S}_N,$$

become a constant permutation independent of $\mathbf{s}^{(n)}$, i.e., $\hat{\mathbf{\Pi}}(\mathbf{s}^{(n)}) = \hat{\mathbf{\Pi}}, \quad \forall \mathbf{s}^{(n)} \in \mathcal{S}_N$. That is,

$$\hat{\mathbf{g}}(\mathbf{x}^{(n)}) = \hat{\mathbf{\Pi}}\hat{\boldsymbol{\rho}}(\mathbf{s}^{(n)}), \quad \forall \mathbf{s}^{(n)} \in \mathcal{S}_N.$$

Lemma C.3. Assume that there is a finite set $\mathcal{S}_N := \{\mathbf{s}^{(1)}, \dots, \mathbf{s}^{(N)}\}$ with $\mathcal{X}_N := \{\mathbf{x} \in \mathcal{X} : \mathbf{x} = \mathbf{f}(\mathbf{s}), \forall \mathbf{s} \in \mathcal{S}_N\}$ such that the Assumption 3.1 is satisfied at each of the N points in $\mathcal{S}_N$. Denote $\hat{\mathbf{g}} \in \mathcal{G}$ as the optimal encoder learned by J-VolMax. Suppose that

$$\hat{\mathbf{g}}(\mathbf{x}^{(n)}) = \hat{\mathbf{\Pi}}(\mathbf{s}^{(n)})\hat{\boldsymbol{\rho}}(\mathbf{s}^{(n)}), \quad \forall \mathbf{s}^{(n)} \in \mathcal{S}_N \quad (98)$$

for a $\mathbf{s}^{(n)}$ -dependent permutation $\hat{\mathbf{\Pi}}$, and additionally these three regularity conditions hold:

1. The functions $\mathbf{f}, \hat{\mathbf{g}}, \hat{\boldsymbol{\rho}}$ are Lipschitz continuous with constants $L_{\mathbf{f}}, L_{\hat{\mathbf{g}}}, L_{\hat{\boldsymbol{\rho}}} > 0$.
2. There is a constant $\gamma > 0$ such that for any permutation matrix $\mathbf{\Pi} \in \mathcal{P}_d$ and $\mathbf{\Pi} \neq \hat{\mathbf{\Pi}}(\mathbf{s}^{(n)})$,

$$\|\hat{\mathbf{g}}(\mathbf{x}^{(n)}) - \mathbf{\Pi}\hat{\boldsymbol{\rho}}(\mathbf{s}^{(n)})\|_2 \geq \gamma, \quad \forall n \in [N]. \quad (99)$$

3. For $\mathcal{N}^{(n)} = \{\mathbf{s} \in \mathcal{S} : \|\mathbf{s} - \mathbf{s}^{(n)}\|_2 < r^{(n)}\}$ with $r^{(n)} < \frac{\gamma}{2(L_{\mathbf{f}}L_{\hat{\mathbf{g}}} + L_{\hat{\boldsymbol{\rho}}})}$, the union set of the neighborhoods, $\mathcal{N} := \bigcup_{n=1}^N \mathcal{N}^{(n)}$, is a connected subset of $\mathcal{S}$.

Then, the permutation ordering $\hat{\mathbf{\Pi}}(\mathbf{s}^{(n)})$ in (98) of the estimated latent components is constant, i.e., $\hat{\mathbf{\Pi}}(\mathbf{s}^{(n)}) = \hat{\mathbf{\Pi}}$ for a fixed permutation $\hat{\mathbf{\Pi}} \in \mathcal{P}_d$. Consequently,

$$\hat{\mathbf{g}}(\mathbf{x}^{(n)}) = \hat{\mathbf{\Pi}}\hat{\boldsymbol{\rho}}(\mathbf{s}^{(n)}), \quad \forall n \in [N]. \quad (100)$$

Proof. Define a function $\ell_{\mathbf{\Pi}}(\mathbf{s})$ parametrized by a permutation matrix $\mathbf{\Pi} \in \mathcal{P}_d$ as

$$\ell_{\mathbf{\Pi}}(\mathbf{s}) := \|\hat{\mathbf{g}}(\mathbf{x}) - \mathbf{\Pi}\hat{\boldsymbol{\rho}}(\mathbf{s})\|_2. \quad (101)$$

From the Lipschitz continuity of $\mathbf{f}, \hat{\mathbf{g}}, \hat{\boldsymbol{\rho}}$, we have the corresponding Lipschitz continuity property of $\ell_{\mathbf{\Pi}}(\cdot)$: for any $\mathbf{s}, \mathbf{s}' \in \mathcal{S}$,

$$|\ell_{\mathbf{\Pi}}(\mathbf{s}) - \ell_{\mathbf{\Pi}}(\mathbf{s}')| = \left| \|\hat{\mathbf{g}}(\mathbf{f}(\mathbf{s})) - \mathbf{\Pi}\hat{\boldsymbol{\rho}}(\mathbf{s})\|_2 - \|\hat{\mathbf{g}}(\mathbf{f}(\mathbf{s}')) - \mathbf{\Pi}\hat{\boldsymbol{\rho}}(\mathbf{s}')\|_2 \right| \quad (102)$$

$$\leq \|(\hat{\mathbf{g}}(\mathbf{f}(\mathbf{s})) - \Pi \hat{\boldsymbol{\rho}}(\mathbf{s})) - (\hat{\mathbf{g}}(\mathbf{f}(\mathbf{s}')) - \Pi \hat{\boldsymbol{\rho}}(\mathbf{s}'))\|_2 \quad (103)$$

$$= \|(\hat{\mathbf{g}}(\mathbf{f}(\mathbf{s})) - \hat{\mathbf{g}}(\mathbf{f}(\mathbf{s}')) + (\Pi \hat{\boldsymbol{\rho}}(\mathbf{s}') - \Pi \hat{\boldsymbol{\rho}}(\mathbf{s})))\|_2 \quad (104)$$

$$\leq \|\hat{\mathbf{g}}(\mathbf{f}(\mathbf{s})) - \hat{\mathbf{g}}(\mathbf{f}(\mathbf{s}'))\|_2 + \|\Pi \hat{\boldsymbol{\rho}}(\mathbf{s}') - \Pi \hat{\boldsymbol{\rho}}(\mathbf{s}')\|_2 \quad (105)$$

$$= \|\hat{\mathbf{g}}(\mathbf{f}(\mathbf{s})) - \hat{\mathbf{g}}(\mathbf{f}(\mathbf{s}'))\|_2 + \|\hat{\boldsymbol{\rho}}(\mathbf{s}') - \hat{\boldsymbol{\rho}}(\mathbf{s}')\|_2 \quad (106)$$

$$\leq L_{\hat{\mathbf{g}}} L_{\mathbf{f}} \|\mathbf{s} - \mathbf{s}'\|_2 + L_{\hat{\boldsymbol{\rho}}} \|\mathbf{s} - \mathbf{s}'\|_2 \quad (107)$$

$$= (L_{\hat{\mathbf{g}}} L_{\mathbf{f}} + L_{\hat{\boldsymbol{\rho}}}) \|\mathbf{s} - \mathbf{s}'\|_2, \quad (108)$$

where we applied the triangle inequality in (103) and (105), the fact that $\Pi \in \mathcal{P}_d$ is an orthogonal matrix in (106), and the Lipschitz continuity of $\mathbf{f}, \hat{\mathbf{g}}, \hat{\boldsymbol{\rho}}$ in (107). Then, we have the following chain of inequalities for any $\mathbf{s} \in \mathcal{N}^{(n)}$:

$$\ell_{\hat{\Pi}(\mathbf{s}^{(n)})}(\mathbf{s}) \leq \ell_{\hat{\Pi}(\mathbf{s}^{(n)})}(\mathbf{s}^{(n)}) + (L_{\hat{\mathbf{g}}} L_{\mathbf{f}} + L_{\hat{\boldsymbol{\rho}}}) \|\mathbf{s} - \mathbf{s}^{(n)}\|_2 \quad (109)$$

$$< \ell_{\hat{\Pi}(\mathbf{s}^{(n)})}(\mathbf{s}^{(n)}) + (L_{\hat{\mathbf{g}}} L_{\mathbf{f}} + L_{\hat{\boldsymbol{\rho}}}) r^{(n)} \quad (110)$$

$$= (L_{\hat{\mathbf{g}}} L_{\mathbf{f}} + L_{\hat{\boldsymbol{\rho}}}) r^{(n)}, \quad (111)$$

where we have used the Lipschitz property of $\ell_{\hat{\Pi}(\mathbf{s}^{(n)})}$ for (109), the defined radius of $\mathcal{N}^{(n)}$ for (110), and the fact that $\ell_{\hat{\Pi}(\mathbf{s}^{(n)})}(\mathbf{s}^{(n)}) = 0$ due to (98) for (111). Similarly, for any $\mathbf{s} \in \mathcal{N}^{(n)}$ and any permutation matrix $\Pi \neq \hat{\Pi}(\mathbf{s}^{(n)})$,

$$\ell_{\Pi}(\mathbf{s}) \geq \ell_{\Pi}(\mathbf{s}^{(n)}) - (L_{\hat{\mathbf{g}}} L_{\mathbf{f}} + L_{\hat{\boldsymbol{\rho}}}) \|\mathbf{s} - \mathbf{s}^{(n)}\|_2 \quad (112)$$

$$> \ell_{\Pi}(\mathbf{s}^{(n)}) - (L_{\hat{\mathbf{g}}} L_{\mathbf{f}} + L_{\hat{\boldsymbol{\rho}}}) r^{(n)} \quad (113)$$

$$> \gamma - (L_{\hat{\mathbf{g}}} L_{\mathbf{f}} + L_{\hat{\boldsymbol{\rho}}}) r^{(n)}, \quad (114)$$

where we again used the Lipschitz property of $\ell_{\hat{\Pi}(\mathbf{s}^{(n)})}$ for (112), the defined radius of $\mathcal{N}^{(n)}$ for (113), and the error gap of non-optimal permutations in assumption (99) for (114). Combining (111) and (114) with the assumption $r^{(n)} < \frac{\gamma}{2(L_{\hat{\mathbf{g}}} L_{\mathbf{f}} + L_{\hat{\boldsymbol{\rho}}})}$, we have

$$\ell_{\hat{\Pi}(\mathbf{s}^{(n)})}(\mathbf{s}) < (L_{\hat{\mathbf{g}}} L_{\mathbf{f}} + L_{\hat{\boldsymbol{\rho}}}) r^{(n)} < \gamma - (L_{\hat{\mathbf{g}}} L_{\mathbf{f}} + L_{\hat{\boldsymbol{\rho}}}) r^{(n)} < \ell_{\Pi}(\mathbf{s}). \quad (115)$$

This implies that for any points $\mathbf{s}$ in the $\mathbf{s}^{(n)}$ -centered neighborhood $\mathbf{s} \in \mathcal{N}^{(n)}$, the optimal permutation for $\ell_{\Pi}(\mathbf{s})$ is still the optimal permutation $\hat{\Pi}(\mathbf{s}^{(n)})$ at the center point $\mathbf{s}^{(n)}$; that is,

$$\hat{\Pi}(\mathbf{s}^{(n)}) = \arg \min_{\Pi \in \mathcal{P}_d} \|\hat{\mathbf{g}}(\mathbf{x}) - \Pi \hat{\boldsymbol{\rho}}(\mathbf{s})\|_2, \quad \forall \mathbf{s} \in \mathcal{N}^{(n)}. \quad (116)$$

As a result, $\hat{\Pi}(\cdot) : \mathcal{N}_N \mapsto \mathcal{P}_d$, which maps a latent vector $\mathbf{s} \in \mathcal{N}$ to a permutation matrix in $\mathcal{P}_d$, is a locally constant mapping over the union set $\mathcal{N}$. Since the locally constant $\hat{\Pi}(\cdot)$ maps to a discrete space $\mathcal{P}_d$, the map $\hat{\Pi}(\cdot)$ is indeed a continuous map. Combining the continuity of $\hat{\Pi}(\cdot)$, a map from $\mathcal{N}$ into a discrete space $\mathcal{P}_d$, with the fact that $\mathcal{N}$ is connected, we can conclude that $\hat{\Pi}(\cdot)$ is indeed constant. Hence, $\hat{\Pi}(\mathbf{s}) = \hat{\Pi}$ for any $\mathbf{s} \in \mathcal{N}$, which also implies $\hat{\Pi}(\mathbf{s}^{(n)}) = \hat{\Pi}$ for any $n \in [N]$.

In conclusion, there is a single permutation matrix $\hat{\Pi} \in \mathcal{P}_d$ such that

$$\hat{\mathbf{g}}(\mathbf{x}^{(n)}) = \hat{\Pi} \hat{\boldsymbol{\rho}}(\mathbf{s}^{(n)}), \quad \forall n \in [N]. \quad (117)$$

□

Lastly, we combine the results from Lemma C.1, C.2, C.3 with a Rademacher complexity-based generalization bound to derive Theorem 3.3.

Theorem 3.3 (Identifiability under Finite-sample SDI). *Assume that there is a finite set $\mathcal{S}_N := \{\mathbf{s}^{(1)}, \dots, \mathbf{s}^{(N)}\}$ with $\mathcal{X}_N := \{\mathbf{x} \in \mathcal{X} : \mathbf{x} = \mathbf{f}(\mathbf{s}), \forall \mathbf{s} \in \mathcal{S}_N\}$ such that the Assumption 3.1 is satisfied at each of the N points in $\mathcal{S}_N$. Let $\hat{\mathbf{g}} \in \mathcal{G}$ be the optimal encoder by J-VolMax criterion (7), and Θ contains the parameters of the learned encoder and decoder. Further assume that the following regularity conditions hold:*

1. The functions $\mathbf{g} = \mathbf{f}^{-1}$ and $\mathbf{g}_{\phi}$ are from classes $\mathcal{G}'$ and $\mathcal{G}$, respectively, where $\mathcal{G}' \subseteq \mathcal{G}$.
2. The functions $\mathbf{f}, \hat{\mathbf{g}}, \hat{\boldsymbol{\rho}}$ are Lipschitz continuous with constants $L_{\mathbf{f}}, L_{\hat{\mathbf{g}}}, L_{\hat{\boldsymbol{\rho}}} > 0$.

3. There is a $\gamma > 0$ such that for any permutation matrix $\mathbf{\Pi} \in \mathcal{P}_d$ and $\mathbf{\Pi} \neq \widehat{\mathbf{\Pi}}(\mathbf{s}^{(n)})$ (from (9)),

$$\|\widehat{\mathbf{g}}(\mathbf{x}^{(n)}) - \mathbf{\Pi}\widehat{\boldsymbol{\rho}}(\mathbf{s}^{(n)})\|_2 \geq \gamma, \quad \forall n \in [N].$$

4. For $\mathcal{N}^{(n)} = \{\mathbf{s} \in \mathcal{S} : \|\mathbf{s} - \mathbf{s}^{(n)}\|_2 < r^{(n)}\}$ with $r^{(n)} < \frac{\gamma}{2(L_{\mathbf{f}}L_{\widehat{\mathbf{g}}} + L_{\widehat{\boldsymbol{\rho}}})}$, the union of the neighborhoods, $\mathcal{N} := \bigcup_{n=1}^N \mathcal{N}^{(n)}$, is a connected subset of $\mathcal{S}$ and $V(\mathbf{s}; \overline{\boldsymbol{\Theta}})$ is optimal for any $\mathbf{s} \in \mathcal{N}$.

5. The points $\mathbf{s}^{(1)}, \dots, \mathbf{s}^{(N)} \in \mathcal{S}_N$ densely locate in $\mathcal{S}$ such that

$$\mathbb{P}(\mathbf{s} \in \mathcal{N}) / \mathbb{P}(\mathbf{s} \in \mathcal{S} \setminus \mathcal{N}) > G_{\max} / G_{\min} > 1, \quad (10)$$

where $G_{\min}, G_{\max}$ are bi-Lipschitz constants of the Jacobian volume surrogate: for any parameters $\boldsymbol{\Theta}_1, \boldsymbol{\Theta}_2$ in $(\mathcal{F}, \mathcal{G})$,

$$G_{\min} \|\boldsymbol{\Theta}_1 - \boldsymbol{\Theta}_2\|_2 \leq |V(\mathbf{s}; \boldsymbol{\Theta}_1) - V(\mathbf{s}; \boldsymbol{\Theta}_2)| \leq G_{\max} \|\boldsymbol{\Theta}_1 - \boldsymbol{\Theta}_2\|_2, \quad \forall \mathbf{s} \in \mathcal{S}. \quad (11)$$

Then, $\widehat{\mathbf{g}}(\mathbf{x}^{(n)}) = \widehat{\mathbf{\Pi}}\widehat{\boldsymbol{\rho}}(\mathbf{s}^{(n)})$, $\forall n \in [N]$ for a constant permutation matrix $\widehat{\mathbf{\Pi}} \in \mathcal{P}_d$. Furthermore, with probability at least $1 - \delta$,

$$\mathbb{E}_{\mathbf{s} \sim p(\mathbf{s})} [\|\widehat{\mathbf{g}}(\mathbf{x}) - \widehat{\mathbf{\Pi}}\widehat{\boldsymbol{\rho}}(\mathbf{s})\|_2] = \mathcal{O} \left((L_{\mathbf{f}}L_{\widehat{\mathbf{g}}} + L_{\widehat{\boldsymbol{\rho}}}) \mathcal{R}_N(\mathcal{G}) + \sqrt{\ln(1/\delta)/N} \right), \quad (12)$$

where $\mathcal{R}_N(\mathcal{G})$ is the empirical Rademacher complexity of the encoder class.

Proof. Given that the J-VolMax criterion (7) provides us with an optimal solution $(\widehat{\mathbf{f}}, \widehat{\mathbf{g}})$ reaching maximal Jacobian volume at each $\mathbf{s}^{(n)}$, such that the estimated latent components $\widehat{\mathbf{g}}(\mathbf{x}^{(n)}) = \widehat{\mathbf{s}}^{(n)}$ identifies $\mathbf{s}^{(n)} \in \mathcal{S}_N$, up to permutation and component-wise transformation, i.e.,

$$\widehat{\mathbf{g}}(\mathbf{x}^{(n)}) = \widehat{\mathbf{\Pi}}\widehat{\boldsymbol{\rho}}(\mathbf{s}^{(n)}), \quad \forall n \in [N], \quad (118)$$

we now show a finite-sample analysis on how well $\widehat{\mathbf{g}}$ can estimate the ground-truth latent sources over all $\mathbf{s} \in \mathcal{S}$, up to the said permutation $\widehat{\mathbf{\Pi}}$ and the invertible element-wise mappings $\widehat{\boldsymbol{\rho}}(\cdot)$ associated with $\mathbf{s}^{(n)} \in \mathcal{S}_N$. To proceed, let us define a loss function

$$\ell(\mathbf{g}, (\mathbf{x}, \mathbf{s})) = \|\mathbf{g}(\mathbf{x}) - \widehat{\mathbf{\Pi}}\widehat{\boldsymbol{\rho}}(\mathbf{s})\|_2, \quad (119)$$

which is L_ℓ -Lipschitz, with the constant depending on Lipschitz continuity of $\widehat{\mathbf{g}}, \mathbf{f}, \widehat{\boldsymbol{\rho}}$ as $L_\ell = L_{\widehat{\mathbf{g}}}L_{\mathbf{f}} + L_{\widehat{\boldsymbol{\rho}}}$ (as derived in Lemma C.3). The defined loss function can be assumed to be upper-bounded by a finite constant as $\ell(\mathbf{g}, (\mathbf{x}, \mathbf{s})) \leq M$. Note that the encoder from J-VolMax $\widehat{\mathbf{g}}$ minimizes the following empirical risk:

$$\widehat{\mathcal{L}}(\mathbf{g}) := \frac{1}{N} \sum_{n=1}^N \ell(\mathbf{g}, (\mathbf{x}^{(n)}, \mathbf{s}^{(n)})), \quad (120)$$

i.e., $\widehat{\mathcal{L}}(\widehat{\mathbf{g}}) = 0$. We are now in a position to use a generalization bound that characterizes the difference between the given empirical risk $\widehat{\mathcal{L}}(\phi)$ and the population risk

$$\mathcal{L}(\mathbf{g}) := \mathbb{E}_{\mathbf{s} \sim p_s} [\ell(\mathbf{g}, (\mathbf{x}, \mathbf{s}))] = \mathbb{E}_{\mathbf{s} \sim p_s} [\|\mathbf{g}(\mathbf{x}) - \widehat{\mathbf{\Pi}}\widehat{\boldsymbol{\rho}}(\mathbf{s})\|_2]. \quad (121)$$

Applying an (empirical) Rademacher complexity-based generalization bound [70, Theorem 26.5] with the contraction lemma [70, Lemma 26.9] gives the following with probability at least $1 - \delta$:

$$\mathcal{L}(\mathbf{g}) \leq \widehat{\mathcal{L}}(\mathbf{g}) + 2L_\ell \mathcal{R}_N(\mathcal{G}) + 4M \sqrt{\frac{2 \ln(4/\delta)}{N}}, \quad \forall \mathbf{g} \in \mathcal{G}. \quad (122)$$

As a result, the following bound holds for any optimal solution $\widehat{\mathbf{g}} \in \mathcal{G}$: with probability at least $1 - \delta$,

$$\mathcal{L}(\widehat{\mathbf{g}}) = \mathbb{E}_{\mathbf{s} \sim p_s} [\|\widehat{\mathbf{g}}(\mathbf{x}) - \widehat{\mathbf{\Pi}}\widehat{\boldsymbol{\rho}}(\mathbf{s})\|_2] \leq 2L_\ell \mathcal{R}_N(\mathcal{G}) + 4M \sqrt{\frac{2 \ln(4/\delta)}{N}}. \quad (123)$$

□

D Additional Remarks on Related Works

IMA. The IMA framework also exploits influence diversity [20]. There, the ground-truth decoder $\mathbf{f}$ is assumed to have an orthogonal Jacobian (e.g., Möbius transformations [20]), which reflects linearly uncorrelated influences from the latent components. Similar orthogonal Jacobian-based regularization have found empirical successes for disentanglement problems in computer vision [71]. However, IMA does not provide identifiability (only *local* identifiability was shown [44]).

Sparse Jacobian-based NMML. Our proposed SDI condition naturally subsumes some cases of sparse Jacobian conditions employed in [2, 3, 19, 49]. A notable example is when $m = 2d$, the SDI assumption boils down to a sparsity pattern in $\mathbf{J}_f$ where each row touches a corner of the weighted L_1 ball $\mathcal{B}_1^{w(s)}$ —and thus the gradients are scaled unit vectors (with $d - 1$ zeros). Nonetheless, when $m > 2d$, completely dense $\mathbf{J}_f$ can also satisfy SDI. We note that although the L_1 regularization on $\mathbf{J}_f$ appears in both J-VolMax and sparse Jacobian criteria (see [3, 19, 50]), the reasons of having this regularization are very different: the latter use the L_1 norm as a surrogate to attain Jacobian sparsity, but our method uses the regularization to confine $\nabla f_i(\mathbf{s})$'s in an L_1 -norm ball (which is not a proxy for sparsity).

Object-centric Representation Learning (OCRL). Many OCRL works also share the perspective of modeling influences—particularly the influences of latent variables onto objects/concepts in the observed domain. While the goal of NMML is to recover each individual latent variable, OCRL [2, 23, 24, 48] seeks a weaker form of identifiability—identifying blocks of variables corresponding to observed objects/concepts. When the block size becomes 1, the goal of OCRL becomes latent variable identification as in NMML. Notably, [2, 23, 24] propose specific forms of $\mathbf{f}$ (which is called “decoder” in the OCRL literature) for block-wise identifiability. These structures often imply sparsity in Jacobian or higher-derivatives of $\mathbf{f}$. For example, the *compositional generator* in [2] and *additive decoder* in [23] are associated with structured sparse $\mathbf{J}_f$'s. Similar to DICA, the work [2] uses a Jacobian-based regularizer in their training loss. The work [23] imposes the additive decoder structure at the model architecture level. Interestingly, the additive decoder structure in [23] was shown to have the same disentanglement effects as a block-diagonal Hessian penalty proposed in [51].

E Experiment Details

E.1 Further Details on Implementation and Evaluation

On Warm-up Heuristic. The use of warm-up heuristic with regularization schedulers help alleviate the numerical instability that comes with optimizing log det of Jacobian of the decoding neural network $\mathbf{f}_\theta$, which can quickly explode if not controlled. Specifically, by gradually introducing the log det term, we prioritize optimizing for data reconstruction in the warm-up period, since it is a hard constraint in the J-VolMax formulation (7c). Moreover, by minimizing $\|\mathbf{J}_{f_\theta}(\mathbf{g}_\phi(\mathbf{x}^{(n)}))\|_1$ during warm-up, we prevent the Jacobian from exploding as a result of maximizing $c_{\text{vol}}(t)$ while keeping the norm term $\|\mathbf{J}_{f_\theta}(\mathbf{g}_\phi(\mathbf{x}^{(n)}))\|_1$ at a reasonable magnitude. This also encourages a smoother encoder $\mathbf{f}_\theta$ with small enough Lipschitz constant; recall that this is important for identifiability of J-VolMax under finite SDI-satisfying samples, as pointed out by Theorem 3.3.

Efficient Computation of Jacobian Volume Regularizer. As mentioned in Section 5, the computational cost of the log-det volume surrogate c_{vol} in (14) is $\mathcal{O}(m^3)$. This might make c_{vol} impractical to use in high-dimensional data setting, where the dimension of the Jacobian matrix is large. An alternative volume surrogate is to use [34, 72]

$$\hat{c}_{\text{tr}} = \text{Tr}((d\mathbf{I} - \mathbf{1}\mathbf{1}^\top)\mathbf{J}_{f_\theta}(\mathbf{g}_\phi(\mathbf{x}^{(n)}))^\top \mathbf{J}_{f_\theta}(\mathbf{g}_\phi(\mathbf{x}^{(n)}))) \quad (124)$$

$$= \sum_{i=1}^{d-1} \sum_{j=i+1}^d \left\| \frac{\partial \mathbf{f}_\theta(\mathbf{g}_\phi(\mathbf{x}^{(n)}))}{\partial \hat{s}_i^{(n)}} - \frac{\partial \mathbf{f}_\theta(\mathbf{g}_\phi(\mathbf{x}^{(n)}))}{\partial \hat{s}_j^{(n)}} \right\|_2^2 \quad (125)$$

which corresponds to the sum of Euclidean distances between all pairs of partial derivative vectors $\partial \mathbf{f}_\theta / \partial \hat{s}_i$. Maximizing $\hat{c}_{\text{tr}}$ would force the partial derivatives to be spread out, akin to the effect from maximizing log-det surrogate c_{vol} , while reducing the computational cost to $\mathcal{O}(m^2)$. In Table 3, we test the wall-clock time needed for backpropagation per epoch of the logdet-based and trace-based regularization, using a setting similar to Mixture C in our synthetic experiment (with $d = 3, m = 40$).

Table 3: Compare performance and wall-clock time for gradient computation (per epoch) of trace-based and logdet-based surrogate for Jacobian volume.

	Logdet-based DICA	Trace-based DICA	Sparse	Base
R^2 score	0.94 ± 0.08	0.91 ± 0.07	0.79 ± 0.12	0.63 ± 0.04
Time (ms)	31.31 ± 0.30	24.89 ± 0.26	18.95 ± 0.18	6.90 ± 0.18

One can see that while trace-based DICA formulation is approximately 25% faster than logdet-based DICA, R^2 score only slightly decreases. Lastly, we remark that there are other methods for reducing the computational cost for optimizing Jacobian determinant of a neural network, such as [73].

Evaluation with R^2 . To evaluate the R^2 score, we use nonlinear kernel ridge regression with radial basis function kernel [74], which is an universal function approximator. We train the regression model on a train set, and use the prediction of the trained regressor on a test set to calculate the coefficient of determination. This gives us the nonlinear R^2 score.

Compute Resources. All experiments use one NVIDIA A40 48GB GPU, hosted on a server using Intel Xeon Gold 6148 CPU @ 2.40GHz with 260GB of RAM.

E.2 Synthetic Simulations

For each of the simulation, we generate 30000 samples from the described synthetic data generation processes, of which 90% are for training and 10% are for evaluating the MCC and R^2 scores. We use two ReLU fully-connected neural networks with one 64-neuron hidden layer for the encoder and the decoder. The autoencoder is trained via Adam method [75] with learning rate 10^{-4} for 200 iterations, among which the first 20 epochs are for warm-up. The regularization hyperparameters are chosen by validating from $\{10^{-2}, 10^{-3}, 10^{-4}, 10^{-5}\}$, resulting in $\lambda_{\text{vol}} = 10^{-4}$, $\lambda_{\text{norm}} = 10^{-4}$, $\lambda_{\text{sp}} = 10^{-4}$. The neural networks are initialized via He initialization with uniform distribution [76].

E.3 Single-cell Transcriptomics Analysis

Data Generation with SERGIO Simulator. We extract top 5 TFs that regulate the most number of genes in TRRUST database. Furthermore, we only keep 30% of genes with a single regulator, so as to have a more challenging mixture; this gives us $m = 178$ gene expressions.

To construct the gene regulatory mechanism f for SERGIO generation, we use the extracted high-confident interactions from TRRUST as primary regulating edges with coefficients randomly sampled from $U(1.5, 1.8)$ if the TF is the gene's activator, and from $U(-1.8, -1.5)$ if the TF is the gene's repressor. In addition, to model the potential spurious cross-talks interactions between the TFs and the genes, we add secondary weak regulating edges with coefficients with random signs and their magnitudes uniformly sampled from $U(0.1, 1.0)$.

We simulate 20000 cells of the same cell type, and use the gene expression data of these cells as the observation dataset to train with J-VolMax learning criterion.

Training. We use two ReLU fully-connected neural networks with one 64-neuron hidden layer for encoder and decoder. We use Adam optimizer [75] with learning rate 10^{-4} , and train the autoencoder for 4000 epochs, with the first 1000 epochs for warm-up. The regularization hyperparameters are chosen by validation from $\{10^{-2}, 10^{-3}, 10^{-4}, 10^{-5}\}$ to be $\lambda_{\text{vol}} = 10^{-3}$, $\lambda_{\text{norm}} = 10^{-4}$, $\lambda_{\text{sp}} = 10^{-4}$. The neural networks are initialized using He initialization with normal distribution [76].

F Additional Experiment: Unsupervised Concept Discovery in MNIST

Setting. In this section, we explore further experiments beyond nonlinear mixture model identification tasks, to demonstrate potential applicability of DICA in disentanglement and representation learning. Specifically, we apply DICA to the MNIST dataset [77] to train an autoencoder using J-VolMax loss function. We use convolutional neural networks (CNNs) for both encoder and decoder, with the architectures reported in Table 4. We choose $d = 10$ as the latent dimension, and the observation dimension is $m = 32 \times 32 = 1024$. The regularization hyperparameters are $\lambda_{\text{vol}} = 10^{-4}$, $\lambda_{\text{sp}} = 10^{-4}$, and we optimizing using Adam with learning rate 10^{-3} for 100 iterations, of which the first 50

Table 4: Architectures of encoder and decoder used in MNIST experiment (all use stride 2, pad 1, out_pad 1)

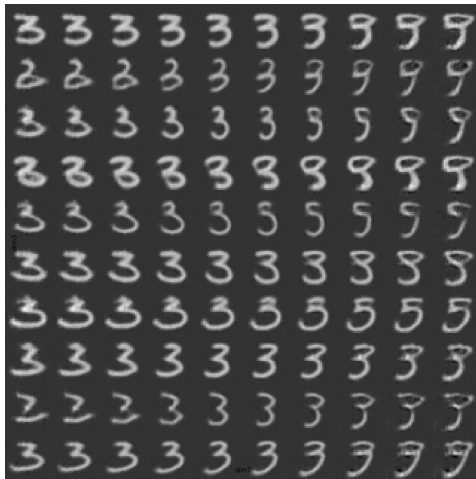
<i>Encoder</i>	<i>Decoder</i>
Input: $x \in \mathbb{R}^{32 \times 32 \times 1}$	Input: $\hat{s} \in \mathbb{R}^{10 \times 1 \times 1}$
3×3 Conv, 256 ReLU	2×2 ConvTrans, 32 ReLU
3×3 Conv, 128 ReLU	3×3 ConvTrans, 64 ReLU
3×3 Conv, 64 ReLU	3×3 ConvTrans, 128 ReLU
3×3 Conv, 32 ReLU	3×3 ConvTrans, 256 ReLU
2×2 Conv, 10	3×3 ConvTrans, 1

epochs are for warm-up. To prevent numerical instability regarding the log-determinant of Jacobian due to high-dimensional data of $m = 1024$, instead of optimizing $\log \det(\cdot)$ directly, we optimize $\log \det(\cdot + \tau \mathbf{I})$ with $\tau > 0$, so as to avoid zero determinant that leads to exploding log-determinant.

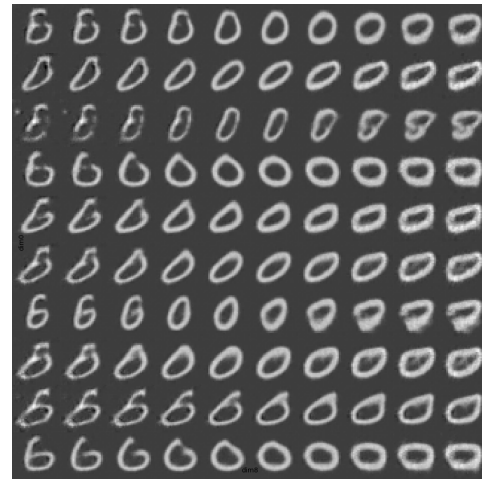
To visualize how the learned latent factors affect the observed images, we encode a test image corresponding to a digit to achieve a latent vector $\hat{s} = [s_1, \dots, s_{10}] \in \mathbb{R}^{10}$. Then, we vary each component s_i in the range of ± 4 std to achieve a new latent vector $\hat{s}_{\text{new}}$ with a certain component being increased/decreased. The new latent vector $\hat{s}_{\text{new}}$ are used as input of the trained decoder to obtain a new image x_{new} . Therefore, we can examine how a specific latent component s_i can affect the observed image x_{new} .

Results. We reported four visualizations corresponding to varying s_8, s_9, s_1, s_7 from images with digit 3, 0, 2, 5, correspondingly in Fig. 3. One can observe that as a learned latent component increases/decreases, the semantic meaning (i.e., digit) of the image changes in a relatively uniform manner towards a new digit. This suggests that the latent components induced by J-VolMax are somewhat correlated with semantic meanings of the images.

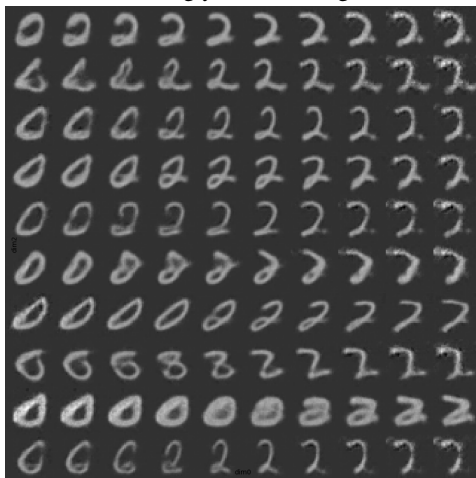
We note that not all latent components we obtained correspond to a clear semantic meaning, and the variation of latent components can sometimes significantly distort the images beyond normal hand-written images. We hypothesize that this is due to the used autoencoder architecture not being suitable for disentanglement purposes, as well as the optimization algorithm not well-designed for this task, since we directly implement J-VolMax criterion without adaptations for image data. Nonetheless, the preliminary results are encouraging, and we speculate that with better-designed architecture and algorithm, the J-VolMax criterion can significantly improve in the challenging task of disentangling meaningful latent components of image data.



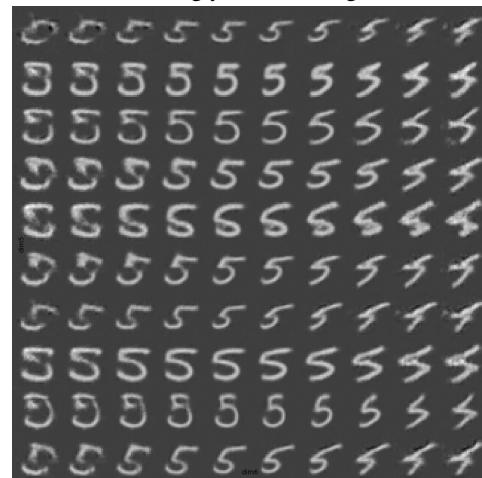
(a) [Anchor digit 3] As s_8 increases, digit 3 increasingly looks like digit 9.



(b) [Anchor digit 0] As s_9 decreases, digit 0 increasingly looks like digit 6.



(c) [Anchor digit 2] As s_1 decreases, digit 2 increasingly looks like digit 0.



(d) [Anchor digit 5] As s_7 increases, digit 5 increasingly looks like digit 4.

Figure 3: Some resulting images obtained by varying a certain component s_i by ± 4 std (increasing from left to right) from the latent vector of an anchor image. Each row corresponds to one of 10 different anchor images sampled from test set, and each column is the resulting image by varying from the corresponding anchor image. We can see that some latent components correlate to the semantic meaning (i.e., digit) of output images: as some s_i increases/decreases, the semantic digit of all 10 anchor images change uniformly towards another digit.

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [\[Yes\]](#)

Justification: Our main claim on model identifiability of DICA framework using J-VolMax learning criterion is proved in Theorem 3.2 and 3.3, and the theoretical results are supported by numerical experiments in Section 5.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: We discussed limitations of the DICA framework, particularly the current algorithm design as well as robustness analysis under noisy setting, in Section 6.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [\[Yes\]](#)

Justification: The assumptions are provided in the statement of Theorem 3.2 and Theorem 3.3, and their proofs are included in the supplementary material.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [\[Yes\]](#)

Justification: An overview of J-VolMax implementation and the datasets used in numerical experiments are provided in Section 5; more details are provided in the supplementary materials.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [\[Yes\]](#)

Justification: The data and code are provided in the supplementary materials, along with relevant instructions.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [\[Yes\]](#)

Justification: The training and test details such as neural network architecture, learning rate, regularization hyperparameters, etc. are provided in the main paper as well as the supplementary materials.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [\[Yes\]](#)

Justification: Experiment results in Section 5.1 include mean and standard deviation, whereas experiment results reported in Section 5.2 are based on median over multiple trials.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.

- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: Compute resources used for the experiments are reported in the supplementary materials.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: The theoretical and experimental works in this paper follow the NeuRIPS Code of Ethics.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [NA]

Justification: The paper is on the theoretical aspects of machine learning, which does not have immediate societal impacts.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: Since this work is of theoretical nature, data and models proposed by this paper does not pose immediate risks for misuse.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: Most of the implementations in this work are created by the authors, and the external libraries, codes, and datasets used in the paper are properly credited, with license and terms respected.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.

- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. **New assets**

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [\[Yes\]](#)

Justification: The codes accompanying this paper are provided in the supplementary materials, together with relevant instructions.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. **Crowdsourcing and research with human subjects**

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [\[NA\]](#)

Justification: This work does not involve crowdsourcing or human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. **Institutional review board (IRB) approvals or equivalent for research with human subjects**

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [\[NA\]](#)

Justification: This paper does not include crowdsourcing or human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.

- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. **Declaration of LLM usage**

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigorousness, or originality of the research, declaration is not required.

Answer: [NA]

Justification: The core method development does not involve LLMs.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.