
Who Guards the Guardians?

The Challenges of Evaluating Identifiability of Learned Representations

Shruti Joshi¹ Théo Saulus¹ Wieland Brendel² Philippe Brouillard¹ Dhanya Sridhar¹ Patrik Reizinger²

¹Mila - Québec AI Institute & Université de Montréal

²Max-Planck-Institute for Intelligent Systems, ELLIS Institute Tübingen, University of Tübingen

Abstract

Identifiability in representation learning is commonly evaluated using standard metrics (e.g., MCC, R^2 , DCI) on synthetic benchmarks with known ground-truth factors. These metrics are assumed to reflect recovery up to the equivalence class guaranteed by identifiability theory. We show that this assumption holds only under specific structural conditions: each metric implicitly encodes assumptions about both the data-generating process (DGP) and the encoder. When these assumptions are violated, metrics become misspecified and can produce systematic false positives and false negatives. Such failures occur both within classical identifiability regimes and in post-hoc settings where identifiability is most needed. We introduce a taxonomy separating DGP assumptions from encoder geometry, use it to characterize the validity domains of existing metrics, and release an evaluation suite for reproducible stress testing and comparison.

1 INTRODUCTION

Learning representations that are interpretable, modular, and controllable is a long-standing goal across machine learning. Identifiability formalises this objective: a representation achieves these properties when it recovers the ground-truth generative factors uniquely, up to a specified equivalence class (Comon, 1994; Hyvärinen and Pajunen, 1999). Strong identifiability guarantees now exist for nonlinear representation learners under auxiliary information (Hyvärinen et al., 2019; Khemakhem et al., 2020a), temporal structure (Hyvärinen and Morioka, 2016), mechanism sparsity (Lachapelle et al., 2022), or for restricted classes of models (Khemakhem et al., 2020c; Marconato et al., 2025). Causal representation learning (CRL) (Schölkopf et al., 2021) builds on these foundations by additionally requiring that the identi-

fied factors admit a causal semantics—typically as variables in a structural causal model with predictable responses under interventions and distribution shifts (Arjovsky et al., 2020; Peters et al., 2016). These results have wide-reaching implications and are increasingly adopted in fields such as mechanistic interpretability (Elhage et al., 2022), where identifiability of learned features is now recognised as a prerequisite for reliable interpretation (Song et al., 2025; Joshi et al., 2025), and in the analysis of pretrained representations more broadly (Roeder et al., 2021).

In practice, these theoretical guarantees are validated empirically. Given ground-truth factors $\mathbf{z} \sim p(\mathbf{z}) \in \mathbb{R}^d$ and learned representation codes $\hat{\mathbf{z}} \in \mathbb{R}^m$, a metric $\mathcal{M}(\mathbf{z}, \hat{\mathbf{z}}) \rightarrow [0, 1]$ returns a scalar interpreted as the degree of identifiability. The standard protocol is to compute $\mathcal{M}$ on a synthetic benchmark with known $\mathbf{z}$, and interpret a high score as evidence that the encoder has recovered the true factors up to a specified equivalence class, e.g., permutation and rescaling.

However, this puts all faith into the metrics—“*Who guards the guardians?*”¹ Each metric encodes structural assumptions about the latent factor distribution $p(\mathbf{z})$, the relationship between ground-truth and learned representation dimensionalities (d and m), the sample size n , and the equivalence class targeted. Yet these assumptions are typically left implicit: papers routinely report a single metric score—MCC (Khemakhem et al., 2020c), R^2 , or DCI-D (Eastwood and Williams, 2018)—as evidence of identifiability, without verifying if the evaluation setting is consistent with the metric’s validity domain. Prior work has observed that metrics can disagree on method rankings and are sensitive to factors such as nonlinearity strength and hyperparameter choice (Seplarskaia et al., 2019; Carbonneau et al., 2022), and that specific metrics produce false positives when latent factors are statistically related (Yao et al., 2025). However, these remain empirical observations tied to particular settings; no prior work characterises *when* and *why* failures arise, nor whether they reflect systematic misspecification predict-

¹“Quis custodiet ipsos custodes?”—Juvenal, *Satires* VI.

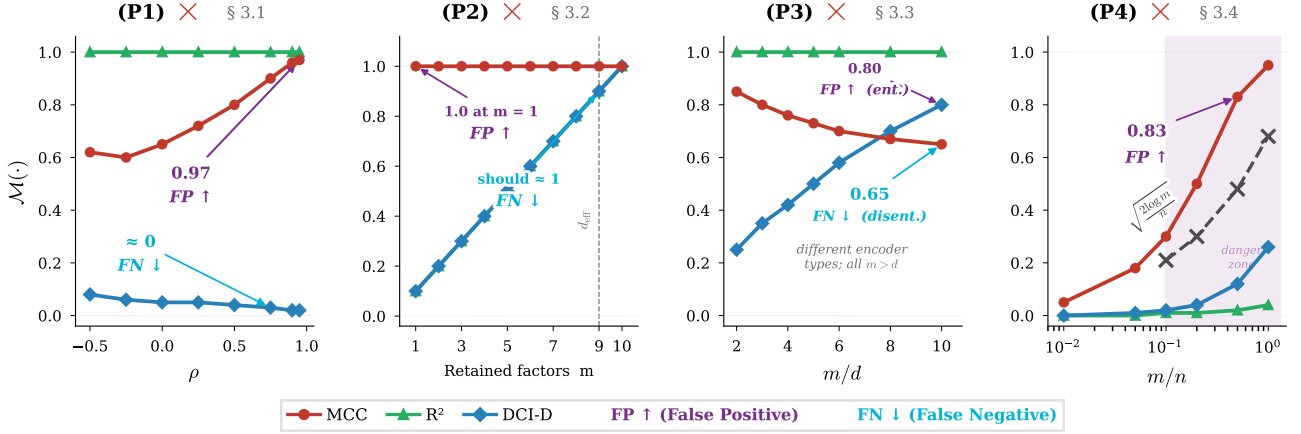


Figure 1: **Every identifiability metric fails under at least one common evaluation setting.** We test four desiderata (Properties 1 to 4) using controlled synthetic encoders that isolate metric behaviour from optimisation artefacts. **(P1)** Latent correlation: MCC conflates correlation with identifiability (FP $\uparrow$ to 0.97); DCI-D penalises it (FN $\downarrow$). **(P2)** Factor dropping: DCI-D reports perfect disentanglement even when 9 of 10 factors are lost. **(P3)** Overcompleteness: MCC inflates for entangled encoders; DCI-D deflates for disentangled ones. **(P4)** Null encoder: all metrics inflate as m/n grows, with MCC scaling in the order of $\sqrt{2 \log m/n}$. R^2 is robust across (P1) and (P4), but shares the (P2) limitation and does not satisfy (P3) across all encoder geometries. No single metric is trustworthy across all settings.

able from each metric’s design. A theorem may guarantee recovery despite correlated factors or only up to an affine transform, whereas, e.g., using MCC targets axis-aligned recovery of independent factors—a strictly stronger assumption whose violation produces systematically wrong scores due to a structural mismatch between what the metric measures and what the experiment intends to measure. This leads to the question:

Can the structural conditions under which a metric faithfully measures identifiability be characterised, and can these conditions be used to predict when false positives and false negatives will arise?

We show that the answer is yes: each metric’s failure modes follow predictably from its encoded assumptions.

Structural misspecification. When a metric’s encoded assumptions do not match the latent factor structure ($p(\mathbf{z})$) or the properties of the encoder producing $\hat{\mathbf{z}}$, we say the metric is *misspecified* for that evaluation setting. Unlike finite-sample noise, misspecification is a population-level property that would persist even when the number of samples $n \rightarrow \infty$, producing *false positives* (high scores despite lack of identifiability) or *false negatives* (low scores despite identifiability up to the desired equivalence class). To predict when and how misspecification arises, we organise assumptions along two orthogonal axes: (i) *latent factor structure*—whether ground-truth factors are independent, correlated, or linked by functional constraints that reduce the effective dimensionality below d ; and (ii) *encoder properties*—the equivalence class, the dimensionality ratio m/d , and if factor information

is distributed across coordinates.

Main contributions. We introduce a two-axis taxonomy (§ 2) separating assumptions about latent factor structure from encoder properties, with formal desiderata for identifiability metrics (Properties 1 to 4). Through controlled synthetic experiments that isolate metric behaviour from optimisation artefacts, we show that no existing metric satisfies all desiderata and characterise precisely how each fails. We derive closed-form analyses showing that (i) MCC approaches 1 when latent factors are highly correlated, even when the encoder remains entangled (§ 3.1), and (ii) the expected MCC under an encoder producing random representations independent of the ground truth is governed by the representation-to-sample ratio (m/n) (§ 3.4). DCI-D is similarly inflated for entangled encoders when $m > d$ (§ 3.3). We also find that a fundamental limitation of all metrics is that they cannot distinguish lossless compression from lossy omission of latent factors when there exist multi-factor dependencies among them (§ 3.2). Detailed discussion of related work appears in § B.

2 A TAXONOMY FOR METRIC (MIS)SPECIFICATION

Identifiable representation learning posits a two-step data generating process.

Formal setup. Ground truth factors $\mathbf{z}$ are sampled first; $\mathbf{z} \sim p(\mathbf{z})$ with $\mathcal{Z} = \text{supp}(p(\mathbf{z})) \subseteq \mathbb{R}^d$. An observation $\mathbf{x} := g(\mathbf{z})$ is then generated via an unknown map $g : \mathcal{Z} \rightarrow \mathcal{X}$ (Hyvärinen and Pajunen, 1999). A learned encoder $f : \mathcal{X} \rightarrow$

$\mathbb{R}^m$ produces $\hat{\mathbf{z}} := f(\mathbf{x})$. It *identifies* the generative factors up to a restricted equivalence class, typically axis-aligned transformations such as permutation and componentwise rescaling, under which the representation $\hat{\mathbf{z}}$ is identified (also called disentangled) (Schmidhuber, 1992; DiCarlo and Cox, 2007; Bengio et al., 2013; Higgins et al., 2018). We adopt the standard notion of identifiability (Hyvärinen and Pajunen, 1999).

Definition 1 (Identifiability up to $\mathcal{G}$). Assume $m = d$ and $\mathcal{G}$ to be a class of transformations acting on $\mathbb{R}^d$. Given some $h \in \mathcal{G}$, where $h : \mathbb{R}^d \rightarrow \mathbb{R}^d$, the encoder f identifies the latent factors up to $\mathcal{G}$ if $f \circ g = h$.

Three standard equivalence classes are: (i) *Permutation and rescaling* ($\mathcal{G}_{\text{perm}}$): $h(\mathbf{z}) = \mathbf{PD}\mathbf{z}$ where $\mathbf{P}$ is a permutation matrix and $\mathbf{D}$ a diagonal scaling matrix, (ii) *Affine* ($\mathcal{G}_{\text{aff}}$): $h(\mathbf{z}) = \mathbf{A}\mathbf{z} + \mathbf{b}$ with $\mathbf{A}$ invertible, and (iii) *Elementwise nonlinear* ($\mathcal{G}_{\text{nl}}$): $h(\mathbf{z}) = (h_1(z_{\pi(1)}), \dots, h_d(z_{\pi(d)}))$ where each h_j is a smooth invertible function. All three make two assumptions: $m = d$, and that each factor is a free degree of freedom, i.e. the *effective dimensionality* d_{eff} (Defn. 2) equals d . When $m \neq d$ (Eastwood et al., 2023; Chen et al., 2025) or $d_{\text{eff}} < d$, Defn. 1 does not apply directly, so we extend it in § 2.2.

Definition 2 (Effective dimensionality). For latent factors $\mathbf{z} \in \mathbb{R}^d$ whose support satisfies k functionally independent smooth constraints $c_1(\mathbf{z}) = 0, \dots, c_k(\mathbf{z}) = 0$ ², the effective dimensionality is $d_{\text{eff}} = d - k$, i.e., the number of factors that can vary independently.

Each metric $\mathcal{M}$ implicitly targets one of the above three equivalence classes, and using a metric outside its target class produces systematically wrong scores. Based on a systematic review of the CRL and nonlinear ICA literature (§ D), we study the three most commonly used metrics: MCC (in two variants: MCC-P based on Pearson correlation, and MCC-S based on Spearman rank correlation), R^2 , and DCI-D (the disentanglement component). MCC computes an optimal one-to-one matching of codes to factors via pairwise correlations, hence targeting elementwise identifiability: (i) in $\mathcal{G}_{\text{perm}}$ (both MCC-P and MCC-S), (ii) and in $\mathcal{G}_{\text{nl}}$ (MCC-S). R^2 is often used by training a linear probe from $\hat{\mathbf{z}}$ to $\mathbf{z}$ and measures explained variance, hence used for evaluating linear identifiability under $\mathcal{G}_{\text{aff}}$ —it cannot distinguish between $\mathcal{G}_{\text{aff}}$ and $\mathcal{G}_{\text{perm}}$. DCI-D trains a probe (linear or nonlinear, e.g., gradient boosted trees (GBT) (Natekin and Knoll, 2013)) to predict each ground-truth factor from the learned codes, then measures how concentrated the resulting feature importances are: a score of 1 means each code is important for predicting at most one factor. Unlike MCC, DCI-D does not require one-to-one code–factor alignment and can handle $m \neq d$, but it remains sensitive to how the probe distributes importance across coefficients—correlated or entangled codes spread

importance across multiple factors, deflating the score even when all information is preserved (§ 3).

To predict metric failures, we consider two orthogonal axes: the *latent factor structure*—whether factors are independent, correlated, or linked by deterministic constraints—and the *encoder geometry*—the equivalence class, dimension ratio m/d , and how factor information is distributed across codes. We define each axis in turn. We use a simple physical system as a running example throughout this section to illustrate how these regimes arise naturally in practice.

RUNNING EXAMPLE: A CIRCUIT WITH A RESISTOR.

Setting. Current flows through a resistor, causing it to dissipate heat and exchange energy with its environment. Sensors (ammeter, voltmeter, thermometer, thermal camera) record the circuit state, producing observations $\mathbf{x}$.

Factors. Four physical quantities influence measurements: T (ambient temperature), R (resistance), I (current), and V (voltage). All four are independently measurable, but not independently variable. Two physical laws constrain them: $R = R_0(1 + \alpha(T - T_0))$ and $V = IR$, where R_0 is the resistance at reference temperature T_0 and α is the material’s temperature coefficient of resistance. So, the factor set (T, R, I, V) has four entries but two degrees of freedom.

Why this matters. An unsupervised learner has no access to Ohm’s law. If it discovers a feature tracking V , that feature is correct even though V is, in principle, determined by I and R . Recovering all four factors reflects the DGP at a particular level of description; a different granularity, say, retaining only (T, I) , is equally valid. Which level is appropriate depends on the downstream task and cannot be determined from the representation alone.

2.1 FACTOR DEPENDENCIES REDUCE EFFECTIVE DIMENSIONALITY

Consider the running example above with four apparent factors: temperature T , current I , resistance R , and voltage V . However, the two physical laws imply that only T and I are free, reducing the effective dimensionality to 2 (instead of 4). Deterministic constraints of this kind are common precisely when ground-truth factors are unknown, such as, when interpreting the representations of a pretrained model, or discovering mechanisms from data—and may arise both from physical laws, or from definitional redundancy (e.g. encoding a position on both a linear and a logarithmic scale) since we do not know which factors should be modelled a priori.

These constraints lie outside the regime in which identifiability theory is usually stated, which assumes each factor is a free degree of freedom ($d_{\text{eff}} = d$). That regime spans inde-

²the Jacobian $[\partial c_i / \partial z_j]$ has rank k

pendently sampled factors, as in standard disentanglement benchmarks (Matthey et al., 2017; Burgess and Kim, 2018; Gondal et al., 2019), and even statistical dependent ones: confounding or noisy causal mechanisms ($z_2 = f(z_1) + \varepsilon, \varepsilon \neq 0$), permitted by several identifiability theorems (Lachapelle et al., 2022; Hyvärinen et al., 2019; Morioka and Hyvärinen, 2024; Khemakhem et al., 2020a,c; Ahuja et al., 2022). We shall see in § 3.1 that this is only in principle, as statistical dependence can already mislead metrics; a *deterministic* constraint ($\varepsilon \equiv 0$) breaks the assumption outright by reducing d_{eff} below d .

Setup. We separate four DGP types, with the first two being standard in the literature that define the regime where identifiability theory operates:

- **$D_{\perp}$ — Independent.** Factors vary independently; each contributes a unique degree of freedom. This is the implicit assumption behind most metrics.
 - ▷ T and I are set by independent exogenous sources.
- **D_{ρ} — Correlated.** Factors are statistically dependent but each retains a unique degree of freedom; no factor is a deterministic function of the others.
 - ▷ A thermostat induces a correlation between T and I .

The next two add a constraint, which removes a degree of freedom ($d_{\text{eff}} = d - k$ for k constraints). Two constraints with the same d_{eff} can still differ in their structure, a distinction the partial information decomposition (Williams and Beer, 2010; Bertschinger et al., 2014) makes precise: a *redundant* constraint, whose dependent factor is predictable from a single other factor, versus a *synergistic* one, whose dependent factor is fixed only by the joint of two or more factors.

- **D_f — Redundant.** The dependent factor is a deterministic, invertible function of a *single* other, $z_j = f(z_i)$. Regressing z_j on z_i leaves zero residual (Pearson $|\text{corr}| = 1$ when f is linear, for a non-linear but monotone f Spearman or mutual information maximal).
 - ▷ $R = R_0(1 + \alpha T)$. Ties 2 factors $\{T, R\}$; one constraint ($k = 1$), $d_{\text{eff}} = d - 1$.
- **D_F — Synergistic.** The dependent factor is a deterministic function of *two or more* others *jointly*, $z_k = g(z_i, z_j)$, so that each argument alone is uncorrelated with z_k —in Pearson ($\text{Cov}(z_i, z_k) = \text{Cov}(z_j, z_k) = 0$) and in rank—while (z_i, z_j) together determine it exactly. The product $z_k = z_i z_j$ of independent, zero-mean factors is the canonical case. The dependence is therefore invisible pairwise and recovery requires conditioning on z_i and z_j together.
 - ▷ $V = IR$. Ties 3 factors $\{I, R, V\}$; one constraint ($k = 1$), $d_{\text{eff}} = d - 1$.

Synergy is not merely a constraint between more than two factors as even though an additive constraint like $z_3 = z_1 + z_2$ relates three factors, it stays pairwise-detectable, since $\text{Cov}(z_1, z_1 + z_2) = \text{Var}(z_1) \neq 0$. D_F instead uses the product $V = IR$, where pairwise correlations between

all variables is zero. Thus d_{eff} (equivalently the constraint count k) sets how much information is dependent, while the redundant/synergistic structure sets whether that dependence is visible to a metric restricted to pairwise statistics or one-to-one matching. Formal description in § E.1.

2.2 ENCODER STRUCTURE AND DIMENSION MISMATCH

Identifiability theory typically assumes $m = d$: the encoder’s output dimension matches the number of latent factors. In practice—particularly when disentangling representations from pretrained models where d is unknown—the common regime is $m > d$, often with $m \gg d$. We organise encoders along two axes: the *equivalence class* up to which factors are identified, and the *dimension ratio* m/d .

Matched dimension. The equivalence classes from Defn. 1 define three encoder types at $m = d$.

- **E1 — Elementwise linear.** Recovery up to $\mathcal{G}_{\text{perm}}$ (Defn. 1 (i)). This is the strongest form of identifiability that can be guaranteed. Ideally, every metric that does not have an intrinsic calibration defect should score 1 here.
 - ▷ $(2T, -R, 0.5I, 3V)$.
- **E2 — Elementwise nonlinear.** Recovery up to $\mathcal{G}_{\text{nl}}$ (Defn. 1 (iii)). Each code is a smooth invertible function of exactly one factor. The parameter α (Tab. 1) controls the degree of nonlinearity; $\alpha = 0$ reduces to E1.
 - ▷ $(\tanh T, R^3, \sqrt[3]{I}, \sinh V)$.
- **E3 — Linearly entangled.** Recovery up to $\mathcal{G}_{\text{aff}}$ (Defn. 1 (ii)). All factor information is preserved, but distributed across coordinates via rotation or shearing. The degree of entanglement is controlled by the condition number κ of $\mathbf{A}$ (Tab. 1): $\kappa = 1$ reduces to E1.
 - ▷ $A(T, R, I, V)^T$: every code mixes all factors.

Dimension mismatch breaks coordinate-wise evaluation.

When $m \neq d$, a single code need not correspond to a single factor. Consider an encoder that represents one factor V by the pair $(\sin V, \cos V)$ such that neither code alone determines V , but the pair identifies V exactly through atan2 . Coordinate-wise notion of identifiability does not credit this, since recovery is only possible through a block of codes rather than through a single code alone. The following definition makes this precise.

Definition 3 (Identifiability under dimension mismatch). *Consider $\emptyset \neq S \subseteq [d]$ to index the recovered factors, and let $\{T_i\}_{i \in S}$ be nonempty, pairwise disjoint code blocks $T_i \subseteq [m]$. The coordinate projections $\pi_S : \mathbb{R}^d \rightarrow \mathbb{R}^{|S|}$ and $\pi_{T_i} : \mathbb{R}^m \rightarrow \mathbb{R}^{|T_i|}$ map $\mathbf{z} \mapsto (z_i)_{i \in S}$ and $\hat{\mathbf{z}} \mapsto (\hat{z}_j)_{j \in T_i}$ respectively. Assume $\mathcal{G}_{|S|}$ to be a transformation class acting on $\mathbb{R}^{|S|}$ (the analogue of the class $\mathcal{G}$ in Defn. 1).*

The encoder $f : X \rightarrow \mathbb{R}^m$ identifies the factors $\mathbf{z}_S$ up to $\mathcal{G}^{|S|}$ through $\{T_i\}_{i \in S}$ if there exist $h \in \mathcal{G}^{|S|}$ and a block-

diagonal readout $\Phi : \mathbb{R}^{|S|} \rightarrow \mathbb{R}^{|T|}$, with each $\varphi_i : \mathbb{R} \rightarrow \mathbb{R}^{|T_i|}$ injective on $h_i(\pi_S(\mathcal{Z}))$ (so each output block reads a single coordinate) mapping $\mathcal{Z}$ to $\mathbb{R}^{|T|}$:

$$\pi_T \circ f \circ g = \Phi \circ h \circ \pi_S \quad \text{on } \mathcal{Z}.$$

Read componentwise, the condition says that code block T_i is a fixed injective readout φ_i of a single recovered factor $h_i \circ \pi_S$, where the transformation $h \in \mathcal{G}^{|S|}$ unmixes the factors while φ_i merely spreads that one recovered factor across the $|T_i|$ codes of its block. Injectivity of φ_i guarantees a readout φ_i^{-1} that collapses the block back to the factor and since both compositions land in $\mathbb{R}^{|T_i|}$, the block size $|T_i|$ is exactly the number of codes devoted to factor i .

For the example above, $\varphi(V) = (\sin V, \cos V)$ with $|T_i| = 2$, $h_i = \text{id}$, and $\varphi^{-1}(a_1, a_2) = \text{atan2}(a_1, a_2)$. Setting $|T_i| = 1$ and $\varphi_i = \text{id}$ for all i recovers the familiar coordinate-wise picture, in which identification means selecting one code per factor; taking $m = d$, $S = [d]$, and $T_i = \{i\}$ reduces Defn. 3 to Defn. 1. Each factor in S is thus encoded by a separate block of codes, and the $m - |\bigcup_i T_i|$ codes outside the blocks are ignored. Next, we can define dimensionality-mismatched encoders.

- **E4 — Undercomplete.** The encoder outputs fewer dimensions than there are ground-truth factors ($m < d$). Since the blocks are disjoint, $|S| \leq m < d$: some factors are unrecoverable for any choice of blocks. While lossy in the standard sense, this may still be a lossless compression when the ground-truth factors are redundant: under $\mathbf{D}_f$ or $\mathbf{D}_F$, an encoder that recovers all d_{eff} independently varying factors (Defn. 2) captures the full information of $\mathbf{z}$. Defn. 3 reports only *which* factors appear in S ; judging whether $|S| \geq d_{\text{eff}}$ constitutes lossless recovery requires, in addition, the constraint structure of the DGP. No current metric makes this distinction: all treat $|S| < d$ uniformly, whether the omitted factors are redundant or independently informative.

▷ $(T, I): |S| = 2 = d_{\text{eff}}$.

- **E5 — Duplicated** ($m > d$, linear). Each factor is copied into one or more codes, every copy a signed rescaling of it. The code–factor map stays one-to-one within each block ($|T_i| = 1$), so any single copy suffices. This is the overcomplete analogue of E1.

▷ $(T, 2T, R, I, V): T$ is encoded twice.

- **E6 — Nonlinear duplicated** ($m > d$, nonlinear). Each code is a smooth invertible function of exactly one factor, and factors can be encoded by multiple codes. No code mixes factors ($|T_i| = 1$, one code per block). This is the overcomplete analogue of E2—the nonlinear counterpart of E5, and the case that separates metrics assuming linear code–factor relations from those invariant to monotone reparametrisation.

▷ $(\tanh T, T^3, R^3, \sqrt[3]{I}, \sinh V): T$ is encoded through two nonlinearities.

- **E7 — Superposed** ($m > d$, linear). Every code is a dense linear mix of all factors, as in E3 but overcomplete.

▷ $\mathbf{A}(T, R, I, V)^T$ with $\mathbf{A} \in \mathbb{R}^{m \times d}$, $\text{rank } \mathbf{A} = d$.

- **E8 — Distributed** ($m > d$, nonlinear). A factor is encoded by a block of several codes ($|T_i| > 1$): the code–factor map is *many-to-one* and no single code recovers the factor. Recovery requires aggregating the block through the readout φ_i^{-1} , whereas coordinate-wise metrics implicitly assume each factor is encoded by a single code ($|T_i| = 1$).

▷ (a_1, a_2, T, R, I) with $(a_1, a_2) = (\sin V, \cos V)$, $V \in (-\pi, \pi]$: the pair jointly encodes V , recovered as $\text{atan2}(a_1, a_2)$, yet neither code alone suffices.

E8 is where Defn. 3 is needed in full generality. For any single code the correspondence is not one-to-one in either direction— $\sin V$ takes the same value at V and $\pi - V$, and $\cos V$ at $\pm V$, so no code determines V and V determines no code invertibly—yet the pair $(\sin V, \cos V)$ determines V exactly. Here identifiability is a property of the block ($|T_i| = 2$, $\varphi_i^{-1} = \text{atan2}$), not of any code, whereas in E5, by contrast, each copy alone suffices ($|T_i| = 1$). A metric that reads codes one at a time therefore cannot be used to evaluate recovery of a distributed factor without a nonlinear, block-level readout (§ 3.3). The distributed factor is the encoder-side dual of a synergistic constraint ($\mathbf{D}_F$): in both, a factor is pinned down only by a *joint* of several variables and never by any one of them alone—those variables being the *codes* of a block for the distributed encoder, and the source *factors* for the synergistic DGP.

E9 and E10 — Control baselines. Respectively $\hat{\mathbf{z}} \sim \mathcal{N}(0, I_m)$ and $\hat{\mathbf{z}} \sim \text{Uniform}([-1, 1]^m)$, independent of $\mathbf{x}$. Every metric should return ≈ 0 . With these definitions in hand, we can formally characterise metric failures. Formal constructions in § E.2.

3 METRICS AS MEASUREMENT INSTRUMENTS

We study the structural sensitivity of identifiability metrics through controlled synthetic experiments. In each experiment, we sample ground-truth factors $\mathbf{z} \in \mathbb{R}^d$ according to a DGP type ($\mathbf{D}_\perp$ – $\mathbf{D}_F$) and construct representations $\hat{\mathbf{z}} = T(\mathbf{z})$ via a transformation matching the encoder type (E1–E10). *The representation encoder is not learned.* This design isolates metric misspecification from optimisation artefacts: every failure we observe is a property of the metric, not of training. Unless otherwise noted, we report results for $n = 1000$ samples, $d = 5$ ground-truth factors, and average over 5 seeds; confidence bands show 95% intervals. For metrics requiring a trained predictor (DCI, R^2), data is split into (80/20) training and test sets. Full experimental details and parameter definitions are in § H³.

³Code for unified implementation of all metrics with improved robustness and our metric evaluation suite.

Table 1: Important parameters and ratios.

Symbol	Meaning	Range
SCALING PARAMETERS		
n	# i.i.d. paired samples	50–10k
m	Dim. of $\hat{\mathbf{z}} \in \mathbb{R}^m$	1–200
d	# ground-truth factors $\mathbf{z} \in \mathbb{R}^d$	2–20
KEY RATIOS		
m/d	Overcompleteness ratio	
d/n	Sample ratio	
m/n	Representation-to-sample ratio	
COMPLEXITY PARAMETERS		
ρ	Pairwise correlation ($\mathbf{D}_\rho$); off-diagonal entries of Σ	$(-1, 1)$
α	Nonlinearity strength (E2); $\alpha = 0$: linear, $[0, 1]$ $\alpha = 1$: fully nonlinear $\hat{z}_j = (1-\alpha)s_j z_{\pi(j)} + \alpha h_j(z_{\pi(j)})$	
κ	Condition # of mixing matrix (E3 , E7); $\kappa = 1$: orthogonal, $\kappa = 50$: ill-cond. $A = U \text{diag}(\text{linspace}(1, \kappa^{-1}, d)) V^\top$	1–50

Metrics evaluated. We evaluate the metrics introduced in § 2, grouped into four families: *correlation-based* (MCC-P/S, MCC-RDC (Lopez-Paz et al., 2013)), *regression-based* (DCI-D, R^2), and *(mutual information) MI-based* (MIG (Chen et al., 2018), InfoMEC (Hsu et al., 2023)), and *conditional independence testing based* (T-MEX (Yao et al., 2025)). In main text, we focus on the commonly used metrics spanning the first two families: MCC, DCI-D, and R^2 .

Sanity checks. We first ask: *do metric scores remain stable when the encoder perfectly recovers each factor (E1), but the DGP varies from independent to correlated to functionally redundant?* Any metric faithful to the equivalence class should return ≈ 1 across $\mathbf{D}_\perp$ – $\mathbf{D}_F$ under E1, since the encoder–factor relationship is identical in all cases. We find that MCC-P, MCC-S, and R^2 have outputs ≈ 1 , while DCI-D exhibits a systematic dip under $\mathbf{D}_F$, particularly at small d (Fig. 7; the dip diminishes as d grows from 5 to 20 but does not vanish). The dip arises because the redundant factor creates collinearity in the regression probe, inflating the importance mass assigned to the dependent factor and reducing the disentanglement score. This persists as n is increased (Fig. 9). A second test is to assess sensitivity to *encoder nonlinearity* rather than DGP structure; Fig. 8 shows flat curves with MCC-S and DCI-D, as expected.

3.1 CORRELATED AND ENTANGLED LATENT FACTORS LEAD TO BOTH FALSE POSITIVES AND FALSE NEGATIVES

We first study how latent-factor correlation ($\mathbf{D}_\rho$) interacts with metric scores under encoders E1 (perfectly disentangled) and E3 (linearly entangled). The encoder is held

fixed; only the pairwise correlation $\rho \in (-1, 1)$ among ground-truth factors varies. Any change in the metric score is therefore a pure artifact of the latent covariance structure.

Property 1 (Invariance to latent correlation). *For $\mathbf{z} \in \mathbb{R}^d$ with pairwise correlations $\text{Corr}(z_i, z_j) = \rho_{ij}$, fix encoder f . A metric $\mathcal{M}$ is invariant to the latent correlation structure if, for every encoder f , $\mathcal{M}(\hat{\mathbf{z}}, \mathbf{z})$ does not depend on $(\rho_{ij})_{i \neq j}$ and only depends on f .*

Violation of Property 1 means $\mathcal{M}$ conflates representation quality with the covariance structure of the DGP.

Setup: $\mathbf{D}_\rho$ + E1/E3. Consider d ground-truth factors with $\mathbf{z} \sim \mathcal{N}(\mathbf{0}, \Sigma)$, $\Sigma_{ii} = 1$, $\Sigma_{ij} = \rho$ for $i \neq j$. Note that the equicorrelation matrices are positive semidefinite only for $\rho \geq -1/(d-1)$; at $d = 10$ this gives $\rho \gtrsim -0.11$, so strongly negative correlations are infeasible at moderate d . E1 is realised as $\hat{z}_j = s_j z_j$ with $s_j > 0$. For E3, the encoder is a full-rank linear map $\hat{\mathbf{z}} = \mathbf{A}\mathbf{z} + \mathbf{b}$ with $\mathbf{A} = \mathbf{U} \text{diag}(\text{linspace}(1, \kappa^{-1}, d)) \mathbf{V}^\top$, where $\mathbf{U}, \mathbf{V}$ are random orthogonal matrices and $\kappa \geq 1$ controls the condition number (degree of entanglement).

Theoretical analysis. We derive a closed-form expression for MCC-P under $\mathbf{D}_\rho$ + E3 (§ G.1), yielding:

Proposition 1 (MCC produces false positives under correlation). *Under $\mathbf{D}_\rho$ + E3, MCC-P depends explicitly on ρ , violating Property 1. Moreover, at both extremes $\rho \rightarrow +1$ and $\rho \rightarrow -1$, $\text{MCC}(\rho) \rightarrow 1$, despite an entangled encoder.*

Prop. 1 predicts not merely a sensitivity issue w.r.t. ρ , but a failure where the metric saturates at 1 even for an entangled encoder identified only up to $\mathcal{G}_{\text{aff}}$. Whereas MCC is designed to distinguish such encoders from ones identified up to $\mathcal{G}_{\text{perm}}$, making an entangled representation indistinguishable from a disentangled one. Under correlated factors and non-axis-aligned encoders, MCC systematically overestimates identifiability, the gap between E1 and E3 narrows as ρ increases (Fig. 2). We observe that the gap and the bias sharpens with growing d . Fig. 11 studies the interaction between ρ and κ at $d = 10$, confirming variation of each metric’s values with ρ rather than κ .

Takeaway. MCC cannot reliably compare representations learned from correlated data, scoring near 1 at high ρ (false positive) for an entangled encoder. DCI-D is overly sensitive to κ , scoring near zero for any non-trivial entanglement (false negative; Fig. 11).

3.2 METRICS CANNOT DETECT MULTI-FACTOR REDUNDANCY

We now study what happens when the encoder outputs fewer dimensions than the number of ground-truth factors ($m < d$). We construct E4 by selecting m factors and applying

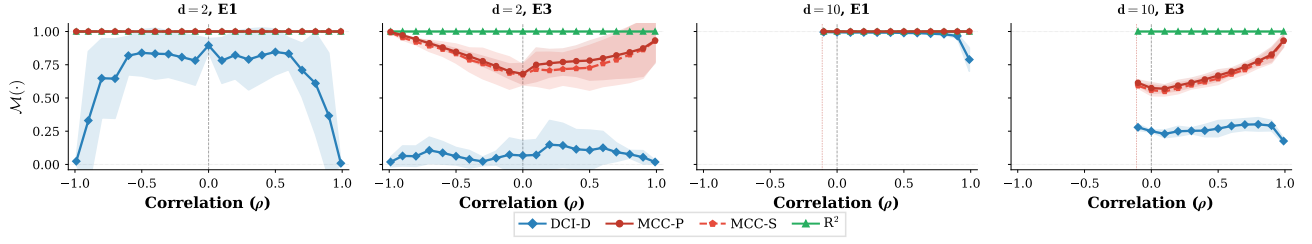


Figure 2: **MCC conflates correlation with identifiability.** *Ideal:* Property 1 implies the metric value should stay parallel to the x -axis. Under E3, MCC increases with ρ and approaches the score of the perfectly disentangled encoder E1 at high correlation, despite the encoder remaining entangled. DCI better separates E1 from E3 but collapses to near-zero scores making it hard to distinguish from a non-identifiable encoder. The bias sharpens with increasing d . See Fig. 13 for all metrics.

elementwise rescaling, so the retained factors are *perfectly* identified. The central question is: *can metrics distinguish an encoder that drops a redundant factor (lossless) from one that drops an informative factor (lossy)?*

Property 2 (Faithfulness to effective dimensionality). *Let $\mathbf{z} \in \mathbb{R}^d$ have effective dimensionality $d_{\text{eff}} \leq d$ (Defn. 2). $\mathcal{M}$ is faithful to the effective dimensionality if $\mathcal{M} = 1$ whenever the encoder recovers all d_{eff} independently varying factors (even if $m < d$), and $\mathcal{M} < 1$ whenever the encoder fails to recover at least one independently varying factor.*

Setup: $\mathbf{D}_\perp/\mathbf{D}_f + \mathbf{E}_4$. Under $\mathbf{D}_\perp$, all d factors are independent, so every omission is lossy. Under $\mathbf{D}_f$, one factor is a deterministic function of another ($z_2 = z_1^3$, so $d_{\text{eff}} = d - 1$); dropping z_2 is lossless. Under $\mathbf{D}_F$, one factor depends on two others ($z_k = g(z_i, z_j)$, $d_{\text{eff}} = 9$); dropping z_k is again lossless. In all cases: $\hat{z}_j = s_j z_j$ for $j \in S$, $|S| = m$.

Fig. 3 reveals a split between metric families. MCC-P/S perform optimal one-to-one matching and score only matched pairs, yielding 1.0 for any $m \geq 1$ regardless of whether omitted factors are redundant or informative. R^2 and DCI-D train a probe to predict *all* d factors from the representation. Under $\mathbf{D}_\perp$ (left), unrepresented factors are unpredictable and $R^2 \approx m/d$. Under $\mathbf{D}_f$ (middle), the redundant factor $z_2 = z_1^3$ is predictable from the retained z_1 , so R^2 and DCI-D stay near 1.0 at $m = d_{\text{eff}}$, thus correctly satisfying Property 2. Under $\mathbf{D}_F$ (correlated factors), $R^2 > m/d$ because the probe partially predicts dropped factors from correlated retained ones; we defer this to Fig. 14.

Under $\mathbf{D}_F$, the redundant factor $z_k = g(z_i, z_j)$ depends jointly on two other factors. Although $d_{\text{eff}} = 9$ (same as $\mathbf{D}_f$), the nonlinear probe fails to detect the relationship Fig. 3 shows a false negative: a lossless encoder is penalised as though it were lossy.

Takeaway. Regression-based metrics (R^2 , DCI-D) detect single-factor redundancy ($\mathbf{D}_f$), but no current metric detects multi-factor redundancy ($\mathbf{D}_F$).

3.3 METRICS CANNOT COMPARE OVERPARAMETRISED ENCODERS

When $m > d$, the encoder outputs more codes than there are factors. We first formalise the desired property we want a metric to exhibit.

Property 3 (Invariance to overcompleteness.). *Let f be an encoder with $m = d$ that identifies factors up to equivalence class $\mathcal{G}$ (Defn. 3), and let f' be an overcomplete encoder ($m > d$) that identifies the same factors up to the same $\mathcal{G}$. A metric $\mathcal{M}$ is invariant to the overcomplete dimension if $|\mathcal{M}(f') - \mathcal{M}(f)| \leq \epsilon(n)$, where $\epsilon(n) \rightarrow 0$ as $n \rightarrow \infty$.*

Violation of Property 3 implies that the metric either spuriously rewards extra codes that add no per-factor information, or that it penalises an encoder that has not lost any factors but merely represents them using multiple codes. In either case, this would represent a metric conflating dimensionality with identifiability.

Setup. We compare four overcomplete geometries (E5–E8) against the matched-dimension entangled baseline E3, under $\mathbf{D}_\perp$. We first fix $m/d = 2$ ($d = 20$, $n = 1600$) and then sweep $m/d \in \{1, 1.5, 2, 3\}$ ($d = 5$, $n = 1000$) to test whether the results are stable as overcompleteness increases. At moderate overcompleteness ($m/d = 2$), all metrics correctly separate entangled from disentangled encoders (Fig. 17). Fig. 4 tests whether this holds as m/d increases.

Fig. 4 shows that increasing m/d does not uniformly increase or decrease scores. Instead, it amplifies the mismatch between each metric’s implicit equivalence class and the encoder’s geometry. Two cases are particularly informative.

MCC cannot be used for distributed codes (E8). Each factor is encoded as k codes (e.g., $\sin z_j, \cos z_j$ for $k = 2$); no single code suffices to recover the factor. MCC pairs each factor with exactly one code, so the best match (say $\sin z_j$) has correlation strictly less than 1 with z_j . As k grows, per-code information thins and MCC-P drops from ~ 0.85 at $m/d = 2$ to ~ 0.65 at $m/d = 10$, even though the factors are fully recoverable from their code subsets. This is a structural

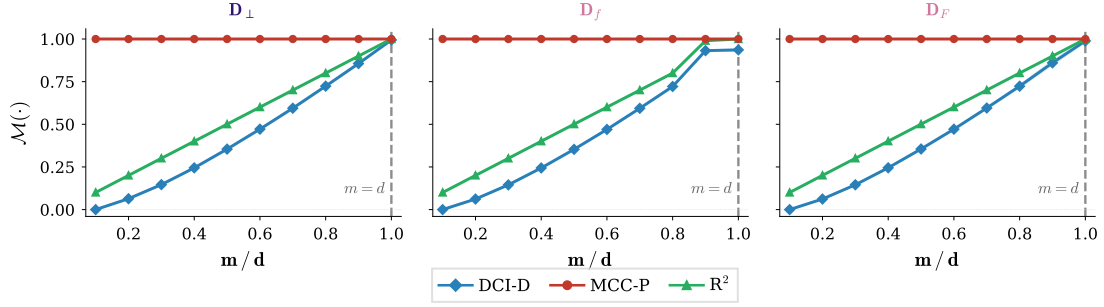


Figure 3: *Ideal*: Property 2 implies that a metric should hold at 1.0 while only redundant factors are dropped ($m \geq d_{\text{eff}}$, lossless compression) and fall only once informative factors are removed ($m < d_{\text{eff}}$). Left ($\mathbf{D}_{\perp}$, $d_{\text{eff}} = d$): every dropped factor is informative. Middle ($\mathbf{D}_f$, $d_{\text{eff}} = d - 1$): the first dropped factor is a *monotone* redundancy ($z_2 = z_1^3$), so removing it ($m/d = 0.9$, $m = d_{\text{eff}}$) is lossless. Right ($\mathbf{D}_F$, $d_{\text{eff}} = d - 1$): the redundancy is *synergistic*. MCC-P **always** stays at 1.0 across all three panels and never register a dropped factor, exhibiting false positives.

failure and it worsens monotonically with m/d . DCI-D does not exhibit this failure as the nonlinear probe can fit all m codes, selecting the k codes in each disjoint subset. Since each selected code predicts only one factor, $D_i = 1$ and DCI-D stays near 1.0 at all tested m/d .

Linear entanglement at high m/d increases DCI-D (E7). However, DCI-D increases substantially even for the linearly entangled encoder, from ~ 0.42 at $m/d = 1.5$ to ~ 0.80 at $m/d = 10$. This produces a false positive that could mislead model comparison.

Only E5 (elementwise linear duplication) satisfies Property 3 across all metrics and all tested m/d values.

Takeaway. No metric satisfies Property 3 across all encoder types; overcomplete representations require multi-metric evaluation or matched-dimension controls.

3.4 HIGH REPRESENTATION-TO-SAMPLE RATIO INCREASES RISK OF FALSE POSITIVES

Unlike the population-level misspecification studied in § 3.1 to 3.3, the false-positive inflation in this section is a finite-sample phenomenon: the bias vanishes as $n \rightarrow \infty$. We include it because the sample regimes encountered in practice—particularly in mechanistic interpretability, where m/n routinely exceeds 1—are far from this asymptotic limit, making the finite-sample floor operationally indistinguishable from structural misspecification.

Property 4 (Insensitivity to uninformative encoders). *For any encoder f independent of $\mathbf{z}$, a metric $\mathcal{M}$ should satisfy $\mathcal{M}(f) \approx 0$ regardless of the dimensionality ratio m/d and sample size n .*

Setup. We construct a null encoder E9 ($\hat{\mathbf{z}} \sim \text{Uniform}([0, 1]^m)$) and sweep over both m/d and m/n .

A metric should assign ≈ 0 to this encoder since it carries no information about $\mathbf{z}$.

MCC violates Property 4 whenever $m/n \geq 0.1$. Reading along any row of Fig. 5 (fixed m/d , varying m/n), scores increase steadily; reading along any column (fixed m/n , varying m/d), scores are approximately constant. The false-positive rate is therefore governed by m/n , not m/d . At $m/n = 0.5$ and $m/d = 1$, MCC-P reports 0.83 for a representation that is pure noise. DCI-D satisfies Property 4 in the large-sample regime, but shows moderate inflation at higher estimation ratios, particularly when m/d is small. R^2 satisfies the property across the entire $(m/d, m/n)$ grid.

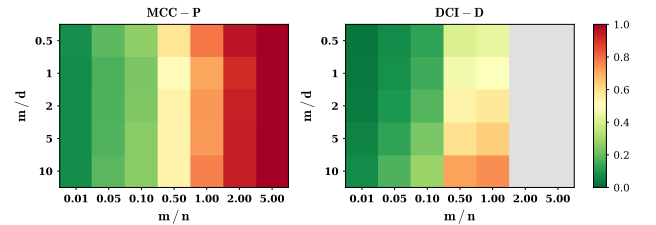


Figure 5: **Null-encoder scores reveal that m/n (columns), not m/d (rows), governs false-positive inflation.** Each cell shows the metric score of a random encoder E9 that carries no information about $\mathbf{z}$; any score above 0 is a false positive. See § G.3 for the theoretical analysis and Fig. 33 for the Gaussian null E10 (nearly identical).

Theoretical analysis. We derive this behaviour in § G.3. Under the null, each entry of the $m \times d$ sample correlation matrix has mean zero and standard deviation $\approx 1/\sqrt{n}$ by the Central Limit Theorem (CLT). Hungarian matching picks the best one-to-one assignment from m candidates per column; the expected maximum of m draws from $\mathcal{N}(0, 1/n)$ scales as $\sqrt{2 \log m/n}$ (Cai and Jiang, 2011), giving $\mathbb{E}[\text{MCC-P}] \gtrsim \sqrt{2 \log m/n}$ (up to a constant). This depends on m and n but not on d , explaining the column-varying, but constant across

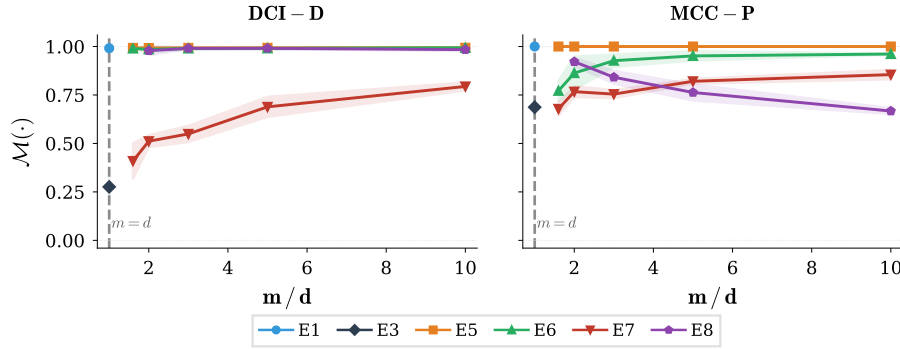


Figure 4: **Sweeping m/d reveals encoder-specific violations of Property 3.** **E5** (elementwise linear duplication) is the only encoder for which all metrics remain stable. DCI-D increases for entangled **E7** as m/d grows. R^2 collapses for nonlinear **E6**. MCC-P decreases for disjoint **E8**. $d = 5$, $n = 1000$.

rows pattern in Fig. 5.

Practical implications. The $m/n \gtrsim 0.1$ threshold is routinely exceeded: evaluating a pretrained LLM such as Llama-3.2-8B ($m = 4096$) with a few hundred samples gives $m/n \in [0.5, 10]$; even standard disentanglement benchmarks with $m = 64$ and $n = 500$ labelled samples yield $m/n > 0.1$. DCI-D requires more samples and exhibits the same inflation. R^2 is the most robust to false positives, but requires $n \gtrsim 500$ under nonlinear encoders (Fig. 9).

Takeaway. MCC is unreliable whenever $m/n \gtrsim 0.1$. Always verify $m \ll n$ and report null-encoder baselines alongside metric scores.

4 CONCLUSION

All existing identifiability metrics can be deceptive (Fig. 1). We provide a taxonomy (§ 2) and theoretical and empirical analyses to characterise these failure modes, then propose four properties (Properties 1 to 4) for future metric design. We distil our findings into a practitioner checklist (§ A) and a metric selection lookup table (Tab. 3). Our results have direct consequences for any pipeline that uses identifiability metrics to make downstream predictions.

Limitations. Our analysis uses synthetic encoders by design, to isolate metric misspecification from optimisation artefacts. The taxonomy does not cover stochastic encoders or discrete factors, all of which arise in practice. Lastly, a systematic study of how metric failures manifest across different families of learned encoders (rather than constructed ones) would be a complementary direction.

ACKNOWLEDGEMENTS

DS acknowledges support from NSERC Discovery Grant RGPIN-2023-04869, and a Canada-CIFAR AI Chair. This

research was enabled in part by computing resources and support provided by Mila-Québec AI Institute.

References

- Kartik Ahuja, Jason S Hartford, and Yoshua Bengio. Weakly supervised representation learning with sparse perturbations. *Advances in Neural Information Processing Systems*, 35:15516–15528, 2022. Cited on page 4.
- Martin Arjovsky, Léon Bottou, Ishaan Gulrajani, and David Lopez-Paz. Invariant risk minimization, 2020. URL <https://arxiv.org/abs/1907.02893>. Cited on page 1.
- Yoshua Bengio, Aaron Courville, and Pascal Vincent. Representation learning: A review and new perspectives. *IEEE transactions on pattern analysis and machine intelligence*, 35(8):1798–1828, 2013. Cited on page 3.
- Nils Bertschinger, Johannes Rauh, Eckehard Olbrich, Jürgen Jost, and Nihat Ay. Quantifying unique information. *Entropy*, 16(4):2161–2183, 2014. Cited on page 4.
- Peter Bühlmann and Sara Van De Geer. *Statistics for high-dimensional data: methods, theory and applications*. Springer Science & Business Media, 2011. Cited on page 26.
- Chris Burgess and Hyunjik Kim. 3d shapes dataset. <https://github.com/deepmind/3dshapes-dataset/>, 2018. Cited on page 4.
- T Tony Cai and Tiefeng Jiang. Limiting laws of coherence of random matrices with applications to testing covariance structure and construction of compressed sensing matrices. *The Annals of Statistics*, 39(3):1496–1525, 2011. Cited on pages 8 and 26.
- Marc-André Carbonneau, Julian Zaidi, Jonathan Boilard, and Ghyslain Gagnon. Measuring disentanglement: A

- review of metrics. *IEEE transactions on neural networks and learning systems*, 35(7):8747–8761, 2022. Cited on pages 1 and 13.
- Guangyi Chen, Yunlong Deng, Peiyuan Zhu, Yan Li, Yifan Shen, Zijian Li, and Kun Zhang. Causalverse: Benchmarking causal representation learning with configurable high-fidelity simulations. In *NeurIPS 2025 Datasets and Benchmarks Track spotlight*, 2025. URL <https://openreview.net/forum?id=R55b2K5KmA>. Cited on page 3.
- Ricky TQ Chen, Xuechen Li, Roger B Grosse, and David K Duvenaud. Isolating sources of disentanglement in variational autoencoders. *Advances in neural information processing systems*, 31, 2018. Cited on pages 6 and 13.
- Pierre Comon. Independent component analysis, a new concept? *Signal processing*, 36(3):287–314, 1994. Cited on pages 1 and 13.
- James J DiCarlo and David D Cox. Untangling invariant object recognition. *Trends in cognitive sciences*, 11(8): 333–341, 2007. Cited on page 3.
- Cian Eastwood and Christopher KI Williams. A framework for the quantitative evaluation of disentangled representations. In *6th International Conference on Learning Representations*, 2018. Cited on pages 1, 13, and 18.
- Cian Eastwood, Andrei Liviu Nicolicioiu, Julius von Kügelgen, Armin Kekic, Frederik Träuble, Andrea Dittadi, and Bernhard Schölkopf. Dci-es: An extended disentanglement framework with connections to identifiability. In *ICLR*, 2023. URL <https://openreview.net/forum?id=462z-gLgSht>. Cited on pages 3 and 14.
- Nelson Elhage, Tristan Hume, Catherine Olsson, Nicholas Schiefer, Tom Henighan, Shauna Kravec, Zac Hatfield-Dodds, Robert Lasenby, Dawn Drain, Carol Chen, et al. Toy models of superposition. *arXiv preprint arXiv:2209.10652*, 2022. Cited on pages 1 and 13.
- Ronald A Fisher. On the mathematical foundations of theoretical statistics. *Philosophical transactions of the Royal Society of London. Series A, containing papers of a mathematical or physical character*, 222(594-604):309–368, 1922. Cited on page 25.
- Muhammad Waleed Gondal, Manuel Wuthrich, Djordje Miladinovic, Francesco Locatello, Martin Breidt, Valentin Volchkov, Joel Akpo, Olivier Bachem, Bernhard Schölkopf, and Stefan Bauer. On the transfer of inductive bias from simulation to the real world: a new disentanglement dataset. *Advances in Neural Information Processing Systems*, 32, 2019. Cited on page 4.
- Irina Higgins, David Amos, David Pfau, Sebastien Racaniere, Loic Matthey, Danilo Rezende, and Alexander Lerchner. Towards a definition of disentangled representations. *arXiv preprint arXiv:1812.02230*, 2018. Cited on page 3.
- Wassily Hoeffding. A class of statistics with asymptotically normal distribution. In *Breakthroughs in statistics: Foundations and basic theory*, pages 308–334. Springer, 1992. Cited on page 25.
- Harold Hotelling and Margaret Richards Pabst. Rank correlation and tests of significance involving no assumption of normality. *The Annals of Mathematical Statistics*, 7(1): 29–43, 1936. Cited on page 26.
- Kyle Hsu, William Dorrell, James Whittington, Jiajun Wu, and Chelsea Finn. Disentanglement via latent quantization. *Advances in Neural Information Processing Systems*, 36: 45463–45488, 2023. Cited on pages 6, 13, and 18.
- Aapo Hyvärinen and Hiroshi Morioka. Unsupervised feature extraction by time-contrastive learning and nonlinear ICA. In *Advances in Neural Information Processing Systems*, 2016. Cited on pages 1 and 13.
- Aapo Hyvärinen and Petteri Pajunen. Nonlinear independent component analysis: Existence and uniqueness results. *Neural networks*, 12(3):429–439, 1999. Cited on pages 1, 2, 3, and 13.
- Aapo Hyvärinen, Hiroshi Sasaki, and Richard E. Turner. Nonlinear ica using auxiliary variables and generalized contrastive learning. *Journal of Machine Learning Research*, 2019. Cited on pages 1, 4, and 13.
- Shruti Joshi, Andrea Dittadi, Sébastien Lachapelle, and Dhanya Sridhar. Identifiable steering via sparse autoencoding of multi-concept shifts, 2025. URL <https://arxiv.org/abs/2502.12179>. Cited on pages 1 and 13.
- Ilyes Khemakhem, Diederik Kingma, Ricardo Monti, and Aapo Hyvärinen. Variational autoencoders and nonlinear ICA: A unifying framework. In *International Conference on Artificial Intelligence and Statistics*, 2020a. Cited on pages 1, 4, and 13.
- Ilyes Khemakhem, Ricardo Monti, Diederik Kingma, and Aapo Hyvärinen. Ice-beem: Identifiable conditional energy-based deep models based on nonlinear ica. *Advances in Neural Information Processing Systems*, 33: 12768–12778, 2020b. Cited on page 13.
- Ilyes Khemakhem, Ricardo Pio Monti, Diederik P. Kingma, and Aapo Hyvärinen. Ice-beem: Identifiable conditional energy-based deep models based on nonlinear ica. In *NeurIPS*, 2020c. URL <https://proceedings.neurips.cc/paper/2020/hash/962e56a8a0b0420d87272a682bfd1e53-Abstract.html>. Cited on pages 1 and 4.

- Sébastien Lachapelle, Pau Rodriguez, Yash Sharma, Katie E Everett, Rémi Le Priol, Alexandre Lacoste, and Simon Lacoste-Julien. Disentanglement via mechanism sparsity regularization: A new principle for nonlinear ica. In *Conference on Causal Learning and Reasoning*, pages 428–484. PMLR, 2022. Cited on pages 1, 4, and 13.
- Francesco Locatello, Stefan Bauer, Mario Lucic, Gunnar Raetsch, Sylvain Gelly, Bernhard Schölkopf, and Olivier Bachem. Challenging common assumptions in the unsupervised learning of disentangled representations. In *International Conference on Machine Learning*, pages 4114–4124. PMLR, 2019. Cited on page 14.
- David Lopez-Paz, Philipp Hennig, and Bernhard Schölkopf. The randomized dependence coefficient. *Advances in neural information processing systems*, 26, 2013. Cited on pages 6 and 18.
- Emanuele Marconato, Sébastien Lachapelle, Sebastian Weichwald, and Luigi Gresele. All or none: Identifiable linear properties of next-token predictors in language modeling. In *AISTATS*, pages 4123–4131, 2025. URL <https://proceedings.mlr.press/v258/marconato25a.html>. Cited on pages 1 and 13.
- Loic Matthey, Irina Higgins, Demis Hassabis, and Alexander Lerchner. dsprites: Disentanglement testing sprites dataset. <https://github.com/deepmind/dsprites-dataset/>, 2017. Cited on page 4.
- Hiroshi Morioka and Aapo Hyvärinen. Causal representation learning made identifiable by grouping of observational variables. In *ICML*, pages 36249–36293, 2024. URL <https://proceedings.mlr.press/v235/morioka24a.html>. Cited on page 4.
- Aaron Mueller, Andrew Lee, Shruti Joshi, Ekdeep Singh Lubana, Dhanya Sridhar, and Patrik Reizinger. From isolation to entanglement: When do interpretability methods identify and disentangle known concepts? In *ACL*, pages 17188–17210, 2026. URL <https://aclanthology.org/2026.acl-long.782/>. Cited on page 13.
- Alexey Natekin and Alois Knoll. Gradient boosting machines, a tutorial. *Frontiers in neurorobotics*, 7:63623, 2013. Cited on page 3.
- Jonas Peters, Peter Bühlmann, and Nicolai Meinshausen. Causal inference by using invariant prediction: identification and confidence intervals. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 78(5):947–1012, 2016. Cited on page 1.
- Geoffrey Roeder, Luke Metz, and Durk Kingma. On linear identifiability of learned representations. In *International Conference on Machine Learning*, pages 9030–9039. PMLR, 2021. Cited on page 1.
- Jürgen Schmidhuber. Learning factorial codes by predictability minimization. *Neural computation*, 4(6):863–879, 1992. Cited on page 3.
- Bernhard Schölkopf, Francesco Locatello, Stefan Bauer, Nan Rosemary Ke, Nal Kalchbrenner, Anirudh Goyal, and Yoshua Bengio. Toward causal representation learning. *Proceedings of the IEEE*, 109(5):612–634, 2021. Cited on pages 1 and 13.
- Anna Sepiarskaia, Julia Kiseleva, and Maarten de Rijke. How to not measure disentanglement. *arXiv preprint arXiv:1910.05587*, 2019. Cited on pages 1 and 13.
- Xiangchen Song, Aashiq Muhamed, Yujia Zheng, Lingjing Kong, Zeyu Tang, Mona T. Diab, Virginia Smith, and Kun Zhang. Position: Mechanistic interpretability should prioritize feature consistency in saes. In *Mech Interp Workshop (NeurIPS 2025) Spotlight*, 2025. URL <https://openreview.net/forum?id=d9ACURK6bI>. Cited on pages 1 and 13.
- P Aditya Sreekar, Ujjwal Tiwari, and Anoop Namboodiri. Mutual information based method for unsupervised disentanglement of video representation. In *2020 25th International Conference on Pattern Recognition (ICPR)*, pages 6396–6403. IEEE, 2021. Cited on page 18.
- Paul L Williams and Randall D Beer. Nonnegative decomposition of multivariate information. *arXiv preprint arXiv:1004.2515*, 2010. Cited on page 4.
- Dingling Yao, Shimeng Huang, Riccardo Cadei, Kun Zhang, and Francesco Locatello. The third pillar of causal analysis? a measurement perspective on causal representations. In *NeurIPS 2025 poster*, 2025. URL <https://openreview.net/forum?id=cx45ACt9Lg>. Cited on pages 1, 6, 13, and 18.

CONTENTS

A	Practitioner Checklist	13
B	Related Work	13
C	Code and Documentation	14
D	Metric Usage Review	14
E	Taxonomy	15
E.1	Data Classes	15
E.2	Encoder Taxonomy	16
F	Metrics	18
G	Expected Metrics' Behaviour: Theory and Derivations	19
G.1	MCC: Correlated latent factors and linear entanglement	19
G.2	Expected behaviour of DCI	24
G.3	MCC false-positive rate under null encoders	25
H	Experiments	27
H.1	Sanity Checks	27
H.2	Correlation among latent factors	30

A PRACTITIONER CHECKLIST

A metric score is interpretable only if two conditions hold: (1) the (DGP, encoder) pair lies in a structurally valid region for that metric, and (2) the sample size n is large enough relative to the relevant dimension to ensure estimation stability. Before reporting scores, verify the following conditions.

Before evaluation.

1. **Check the overparametrisation ratio m/n .** If $m/n > 0.1$, MCC scores are unreliable: the expected score under a null encoder exceeds $\sqrt{2 \log m/n}$ (§ 3.4). Increase n or reduce m before interpreting results.
2. **Report a null-encoder baseline.** Compute every metric on a random or constant encoder with the same (m, n, d) . Without this baseline, false positives are indistinguishable from genuine identifiability (§ 3.4).
3. **Know your DGP assumptions.** Determine whether latent factors are independent ($\mathbf{D}_\perp$) or correlated ($\mathbf{D}_\rho$), and whether the representation is matched ($m = d$), overcomplete ($m > d$), or undercomplete ($m < d$).

Choosing a metric.

4. **Matched dimension, independent factors ($m = d, \mathbf{D}_\perp$):** all three metrics (MCC, DCI-D, R^2) are reliable.
5. **Correlated factors ($\mathbf{D}_\rho$):** prefer R^2 . MCC conflates correlation with identifiability (Prop. 1); DCI-D collapses under moderate entanglement (§ 3.1).
6. **Overcomplete representations ($m > d$):** no single metric is reliable across all encoder geometries. Use multiple metrics and compare against matched-dimension controls (§ 3.3).
7. Consult Tab. 3 for a full lookup table.

Interpreting scores.

8. **A high MCC does not imply identifiability** when m/n is large or factors are correlated.
9. **A high DCI-D does not imply disentanglement** when the encoder is overcomplete and linearly entangled.
10. **No pairwise metric detects multi-factor redundancy ($\mathbf{D}_F$);** higher-order statistics are needed (§ 3.2).

B RELATED WORK

Identifiability theory. Nonlinear ICA (Comon, 1994; Hyvärinen and Pajunen, 1999) establishes sufficient conditions under which latent factors can be recovered up to well-defined equivalence classes. Identifiability guarantees often leverage auxiliary variables (Hyvärinen et al., 2019; Khemakhem et al., 2020a), temporal structure (Hyvärinen and Morioka, 2016), mechanism sparsity (Lachapelle et al., 2022), and restricted model classes (Khemakhem et al., 2020b; Marconato et al., 2025). Causal representation learning extends these results by additionally requiring that identified factors admit causal semantics with predictable behaviour under interventions (Schölkopf et al., 2021). These works establish when identifiability holds in theory. We study whether the metrics used to verify these guarantees empirically are faithful to the equivalence classes the theorems provide. Our results indicate that even with correlations between latent factors, reliability on metrics drops § 3.1.

Identifiability and disentanglement metrics. A substantial body of work has proposed metrics for evaluating learned representations against ground-truth factors, including DCI (Eastwood and Williams, 2018), MIG (Chen et al., 2018), MCC (Khemakhem et al., 2020b), InfoMEC (Hsu et al., 2023), and T-MEX (Yao et al., 2025). Seplarskaia et al. (2019) showed that several metrics disagree on comparing methods and cautioned against relying on a single score. Carbonneau et al. (2022) surveyed metrics and noted the lack of a unified framework connecting metric assumptions to evaluation validity. Our work differs from both. Instead of comparing metric rankings across methods, we identify the structural conditions on the DGP and encoder geometry under which each metric’s score is interpretable, and show that the resulting failure modes are misspecification, not optimisation failures.

Overcomplete representations and mechanistic interpretability. Recent work in mechanistic interpretability uses sparse autoencoders to extract interpretable features from pretrained models (Elhage et al., 2022), and identifiability of these features is increasingly recognised as necessary for reliable interpretation (Song et al., 2025; Joshi et al., 2025; Mueller et al., 2026). These settings are inherently overcomplete $m \gg d$ and sample-constrained $m \gg n$. We show that current metrics are not reliable under overcompleteness: MCC does not work for overcomplete distributed codes, DCI-D may spuriously reward a

linearly entangled representation (§ 3.3), and the high m/n ratios typical of these evaluations may push the metrics into the regime where they can score high even with a random representation (§ 3.4).

Relationship to prior evaluation studies. Locatello et al. (2019) demonstrated that unsupervised disentanglement learning requires inductive biases, studying how *learning algorithms* behave under different model and data assumptions. Our work is complementary: we study how *evaluation metrics* behave under different structural regimes, holding the encoder fixed. Their finding that unsupervised disentanglement is impossible without inductive biases is orthogonal to our finding that even supervised metrics are structurally misspecified under conditions the underlying identifiability theorems explicitly permit. Eastwood et al. (2023) extended DCI to handle dimension mismatch; our Defn. 3 generalises this to arbitrary equivalence classes and connects it to the full DGP taxonomy, revealing failure modes beyond what dimension-mismatch alone predicts.

C CODE AND DOCUMENTATION

All results in this paper are produced by `identifiability-guard`, an open-source Python package that provides a unified, unit-tested implementation of every metric, data-generating process (DGP), and encoder studied here. The source code is available at github.com/shrutij01/identifiability-guard, and browsable documentation is hosted at shrutij01.github.io/identifiability-guard.

Design. The package follows the same three stages as our taxonomy (§ E)—generate ground-truth factors, apply an encoder, then score recovery—so that any DGP can be paired with any encoder and evaluated by any metric through a common interface. It is organised into four modules under `src/identifiability-guard/`: `dgp/` (factor generators), `encoders/` (transformations), `metrics/` (recovery scores), and `evaluation/` (shared utilities), with unit tests in `tests/` and runnable experiment drivers in `examples/`.

Correspondence to the taxonomy. Code objects are named after the labels used throughout the paper, so each symbol maps to one implementation. The DGP classes $\mathbf{D}_1\text{--}\mathbf{D}_F$ (§ E.1) live in `dgp/` (e.g. `D1Independent`); the encoders $\mathbf{E1}\text{--}\mathbf{E8}$ together with the null baselines $\mathbf{E9}/\mathbf{E10}$ (§ E.2) live in `encoders/` as `E1`–`E10` (e.g. `E1ElementwiseLinear`); and the metrics MCC, DCI, R^2 , MIG, T-MEX, and InfoMEC live in `metrics/`, each exposing a common `compute(Z, Z_hat)` entry point (e.g. `MCC`, `DCI`).

Documentation. The documentation site includes a user-guide page for each stage—Data-Generating Processes, Encoder Mixings, and Identifiability Metrics—an auto-generated API reference, worked examples, and installation instructions. It also provides technical notes on numerical stability, computational efficiency, and reproducibility.

Reproducing the experiments. After installation (`pip install -e .`), the scripts in `examples/` regenerate the reported scores: the full $\text{DGP} \times \text{encoder}$ grid is produced by `evaluate_all_combinations_combined.py`, and the sample-size sweeps by `evaluate_sensitivity.py`. Each uses the seeds and sample sizes summarised in Tab. 1, so the figures in § H are reproducible end-to-end.

D METRIC USAGE REVIEW

We conducted a systematic review of evaluation metrics used in causal representation learning (CRL) and nonlinear independent component analysis (ICA). Using the Semantic Scholar API, we retrieved papers published between 2020 and 2025 at major ML conferences (NeurIPS, ICLR, ICML, AISTATS, UAI, AAAI, CLeaR, JMLR) based on the terms ‘causal representation learning’ and ‘nonlinear ICA’. **Among the 62 papers identified, most relied on MCC (25), followed by R^2 (9) and DCI (2).** None employed more recent metrics such as MIG or T-MEX. Finally, several papers did not use standard metrics at all, instead reporting performance in terms of objective optimization or relying on qualitative assessments.

Also, nonlinear ICA papers use MCC (61%) more often than CRL (29%).

Search queries. Semantic Scholar bulk search with two queries — “causal representation learning” and “nonlinear ica” — covering the two dominant terminological traditions, restricted to NeurIPS, ICLR, ICML, AISTATS, UAI, AAAI, CLeaR, and JMLR over 2020–2025.

Inclusion. A paper was kept if it (i) appeared in a listed venue in the specified period and (ii) contained an experimental validation (purely theoretical or position papers excluded).

De-duplication. Result sets merged and de-duplicated by Semantic Scholar paper ID.

Search queries. Semantic Scholar bulk search with two queries — ”causal representation learning” and ”nonlinear ica” — covering the two dominant terminological traditions, restricted to NeurIPS, ICLR, ICML, AISTATS, UAI, AAAI, CLeaR, and JMLR over 2020–2025.

Inclusion. A paper was kept if it (i) appeared in a listed venue in the specified period and (ii) contained an experimental validation (purely theoretical or position papers excluded).

De-duplication. Result sets merged and de-duplicated by Semantic Scholar paper ID.

Evaluation. MCC usage was assessed by one author via manual inspection. We did not compute a formal inter-rater coefficient, but the criterion is a binary, verifiable fact readable from each paper’s results tables, leaving little room for interpretive disagreement. The full annotated 62-paper list will be released with the camera-ready.

E TAXONOMY

E.1 DATA CLASSES

Let $\mathbf{z} = (z_1, \dots, z_d)^\top \in \mathbb{R}^d$ denote the ground-truth latent factors with joint density $p(\mathbf{z})$. We classify the factor distribution along two axes: *statistical dependence* (mutual information) and *functional dependence* (deterministic constraints). Classes $\mathbf{D}_\perp$ – $\mathbf{D}_\rho$ operate within the standard CRL setting ($d_{\text{eff}} = d$); Classes $\mathbf{D}_f$ – $\mathbf{D}_F$ extend it to settings where functional constraints reduce the effective dimensionality ($d_{\text{eff}} < d$; cf. Defn. 2). Plot legends label DGP types by index: $\mathbf{D1} = \mathbf{D}_\perp$ (independent), $\mathbf{D2} = \mathbf{D}_\rho$ (correlated), $\mathbf{D3} = \mathbf{D}_f$ (redundant), $\mathbf{D4} = \mathbf{D}_F$ (synergistic).

$\mathbf{D}_\perp$ — Independent factors. The factors are mutually independent and non-redundant:

$$p(z_1, \dots, z_d) = \prod_{j=1}^d p(z_j), \quad I(z_i; z_j) = 0 \quad \forall i \neq j.$$

No statistical, functional, or structural dependence exists among factors. In particular, $\text{Corr}(z_i, z_j) = 0$ for all $i \neq j$, and $d_{\text{eff}} = d$.

Canonical example: $z_1, z_2 \stackrel{\text{i.i.d.}}{\sim} \text{Uniform}(0, 1)$.

$\mathbf{D}_\rho$ — Correlated (statistically dependent) factors. The factors share information but each retains a unique degree of freedom; no factor is a deterministic function of any subset of the others:

$$\exists i \neq j : I(z_i; z_j) > 0, \quad H(z_j | z_{\setminus j}) > 0 \quad \forall j,$$

where $z_{\setminus j} := (z_1, \dots, z_{j-1}, z_{j+1}, \dots, z_d)$. The first condition asserts statistical dependence; the second asserts non-redundancy: no factor is determined by the rest. Hence $d_{\text{eff}} = d$.

Canonical example: $(z_1, z_2, z_3)^\top \sim \mathcal{N}(0, \Sigma)$, $\Sigma = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & \rho \\ 0 & \rho & 1 \end{pmatrix}$ with $\rho \in (0, 1)$.

Remark. $\mathbf{D}_\rho$ subsumes both causal dependence ($z_j = f(z_i) + \varepsilon$, $\varepsilon \not\equiv 0$) and confounded dependence (shared latent common cause), as well as nonlinear dependence invisible to linear measures. For instance, $z_1 \sim \mathcal{N}(0, 1)$, $z_2 = z_1^2 + \varepsilon$, $\varepsilon \sim \mathcal{N}(0, \sigma^2)$ satisfies $\text{Corr}(z_1, z_2) = 0$ yet $I(z_1; z_2) > 0$: the dependence is real but purely nonlinear. As long as $H(\varepsilon) > 0$, the relationship is non-deterministic and falls under $\mathbf{D}_\rho$. **In this paper, we’ll focus on linear non-deterministic dependence only.**

$\mathbf{D}_f$ — Redundant (single-factor functional constraint). At least one factor is a deterministic function of exactly one other factor, reducing the effective dimensionality:

$$\exists i \neq j, \exists f : \mathbb{R} \rightarrow \mathbb{R} : \quad z_j = f(z_i) \quad \text{a.s.}, \quad H(z_j | z_i) = 0.$$

Two structurally distinct subcases arise:

$\mathbf{D}_f\mathbf{A}$ — Invertible (information-preserving). The map f is injective, so f^{-1} exists. Then $H(z_j) = H(z_i)$ and $I(z_i; z_j) = H(z_i)$. The intrinsic dimension of (z_i, z_j) is 1, but no information is lost. *Canonical example:* $z_2 = z_1^3$.

D_fB — Non-invertible (collapsed). The map f is many-to-one, so f^{-1} does not exist. Then $H(z_j) < H(z_i)$ and $I(z_i; z_j) = H(z_j) < H(z_i)$: information is destroyed. *Canonical example:* $z_2 = \text{sign}(z_1)$.

In both subcases, $d_{\text{eff}} \leq d - 1$ (one constraint removes one degree of freedom).

▸ $R = R_0(1 + \alpha T)$: resistance is an invertible function of temperature (**D_fA**).

In this paper, **D_fA** will be of interest to us.

D_F — Synergistic (multi-factor functional constraint). At least one factor is a deterministic function of two or more other factors, but not of any single one:

$$\exists k, \exists S \subset \{1, \dots, d\} \text{ with } |S| \geq 2 : \quad z_k = g(z_S) \quad \text{a.s.,}$$

where $z_S := (z_j)_{j \in S}$, and no function of a strict subset of z_S determines z_k . Formally:

$$H(z_k | z_S) = 0, \quad H(z_k | z_T) > 0 \quad \forall T \subsetneq S.$$

The constraint cannot be decomposed into single-variable contributions: the dependence is deterministic but *synergistic*. For the canonical zero-mean product below, every pairwise linear correlation vanishes ($\text{Corr}(z_j, z_k) = 0$ for each $j \in S$) even though (z_S) jointly determines z_k , so the dependence is invisible to pairwise statistics.

Canonical example: $z_1 = z_2 z_3$ with $z_2 \perp z_3$ independent and zero-mean.

▸ $V = IR$: voltage is jointly determined by current and resistance. Ties 3 factors $\{I, R, V\}$; one constraint ($k = 1$), $d_{\text{eff}} = d - 1$.

Summary. Under **D₊–D_ρ**, $d_{\text{eff}} = d$ (no functional constraints). Under **D_f–D_F**, $d_{\text{eff}} < d$ (deterministic constraints reduce the number of free degrees of freedom; cf. Defn. 2). The horizontal axis of the validity-domain map (Tab. 3) captures this progression.

E.2 ENCODER TAXONOMY

Let $f : \mathbb{R}^n \rightarrow \mathbb{R}^m$ denote the learned encoder, producing $\hat{\mathbf{z}} := f(\mathbf{x}) = f(g(\mathbf{z})) \in \mathbb{R}^m$. We classify encoders by (i) the equivalence class $\mathcal{G}$ up to which factors are identified (Defn. 1), and (ii) the dimension ratio m/d . Throughout, S_d denotes the symmetric group on $\{1, \dots, d\}$.

Matched dimension ($m = d$).

E1 — Elementwise linear (permutation & rescaling). The encoder identifies each factor up to $\mathcal{G}_{\text{perm}}$:

$$\exists \pi \in S_d, \exists a_j \neq 0 : \quad \hat{z}_j = a_j z_{\pi(j)}, \quad j = 1, \dots, d.$$

No cross-factor mixing or nonlinear reparameterisation is present beyond scaling and permutation. This is the strongest form of identifiability and every metric should score 1.

Canonical example: $\hat{z}_1 = 2z_3$, $\hat{z}_2 = -z_1$, $\hat{z}_3 = 0.5z_2$ (for $d = 3$).

E2 — Elementwise nonlinear (invertible componentwise). The encoder identifies each factor up to $\mathcal{G}_{\text{nl}}$:

$$\exists \pi \in S_d : \quad \hat{z}_j = g_j(z_{\pi(j)}), \quad j = 1, \dots, d,$$

where each $g_j : \mathbb{R} \rightarrow \mathbb{R}$ is a smooth, invertible scalar function. Information is preserved factor-wise, but linear correlation between $\hat{z}_j$ and $z_{\pi(j)}$ may be misleading. The parameter α (Tab. 1) controls the degree of nonlinearity; $\alpha = 0$ reduces to **E1**.

Canonical example: $\hat{z}_j = z_{\pi(j)}^3$.

In practice, we use the following nonlinear functions: $z \mapsto \tanh(2z^3 - 0.1)$, $z \mapsto \sinh(2z^3 + 0.1)$, $z \mapsto \text{sign}(z)|z|^{1/2}$, $z \mapsto z^3$, $z \mapsto z^5$, $z \mapsto \exp(z)$, $z \mapsto \text{sign}(z) \log(1 + |3z + 1|)$.

E3 — Linearly entangled. The encoder identifies factors up to $\mathcal{G}_{\text{aff}}$:

$$\hat{\mathbf{z}} = \mathbf{A} \mathbf{z}, \quad \mathbf{A} \in \mathbb{R}^{d \times d}, \det(\mathbf{A}) \neq 0,$$

with $\mathbf{A}$ not a signed permutation matrix (i.e., at least one row has two or more nonzero entries). All factor information is preserved globally, but individual factors are distributed across coordinates. The condition number $\kappa(\mathbf{A})$ controls the degree of entanglement; $\kappa = 1$ reduces to **E1**.

Canonical example: $\hat{z}_2 = a z_2 + b z_3$ with $ab \neq 0$.

Dimension mismatch ($m \neq d$).

The standard definition of identifiability (Defn. 1) assumes $m = d$. When $m \neq d$, we use the generalised notion of Defn. 3: the encoder identifies a subset $S \subseteq \{1, \dots, d\}$ of factors via a readout $r : \mathbb{R}^m \rightarrow \mathbb{R}^{|S|}$.

E4 — Undercomplete ($m < d$). The encoder outputs fewer dimensions than there are ground-truth factors, so $|S| < d$: some factors are unrecoverable regardless of the readout r . Each retained factor is encoded elementwise:

$$\hat{z}_j = a_j z_{i(j)}, \quad j = 1, \dots, m,$$

with all $i(j)$ distinct and $a_j \neq 0$.

Under **D_f–D_F**, this need not be lossy in the information-theoretic sense: if the encoder recovers all d_{eff} independently varying factors, it captures the full information of $\mathbf{z}$ (Defn. 2). No current metric distinguishes omission of a redundant factor from omission of an informative one.

▷ $(T, I): |S| = 2 = d_{\text{eff}}$.

Overcomplete encoders ($m > d$).

We now define four overcomplete encoder types ($m > d$), each corresponding to a distinct code–factor geometry.

E5 — Duplicated ($m > d$, **linear**). Each output coordinate is a signed rescaling of exactly one ground-truth factor, and each factor is copied into one or more coordinates. Let $\sigma : \{1, \dots, m\} \rightarrow \{1, \dots, d\}$ be a surjective assignment (every factor is represented at least once; some are duplicated), and let $a_j \neq 0$. Then:

$$\hat{z}_j = a_j z_{\sigma(j)}, \quad j = 1, \dots, m, \quad m > d.$$

The surjectivity of σ ensures no factor is lost. Because every copy is an invertible rescaling, the code–factor map is **one-to-one** within each block ($|T_i| = 1$): any single copy already recovers its factor, and the remaining copies are redundant. This is the overcomplete analogue of **E1**.

Example ($d = 2, m = 4$): $\hat{z}_1 = 1.3 z_1$, $\hat{z}_2 = -0.7 z_2$, $\hat{z}_3 = 0.9 z_1$, $\hat{z}_4 = -1.5 z_2$.

E6 — Nonlinear duplicated ($m > d$, **nonlinear**). Each output coordinate is a smooth invertible function of exactly one ground-truth factor, and each factor may be encoded by more than one coordinate. Let $\sigma : \{1, \dots, m\} \rightarrow \{1, \dots, d\}$ be a surjective assignment and each $g_j : \mathbb{R} \rightarrow \mathbb{R}$ an invertible scalar function:

$$\hat{z}_j = g_j(z_{\sigma(j)}), \quad j = 1, \dots, m, \quad m > d.$$

No coordinate mixes factors, so factor-wise information is preserved up to $\mathcal{G}_{\text{nl}}$ and the code–factor map is **one-to-one** within each block ($|T_i| = 1$): any single code recovers its factor. This is the overcomplete analogue of **E2** and the nonlinear counterpart of **E5**.

Example ($d = 2, m = 3$): $\hat{z}_1 = \tanh(z_1)$, $\hat{z}_2 = z_2^3$, $\hat{z}_3 = z_1^3$.

E7 — Superposed ($m > d$, **linear**). The encoder is a dense linear map with $m > d$, the overcomplete analogue of **E3**:

$$\hat{\mathbf{z}} = \mathbf{A} \mathbf{z}, \quad \mathbf{A} \in \mathbb{R}^{m \times d}, \text{rank}(\mathbf{A}) = d,$$

where at least one row of $\mathbf{A}$ has two or more nonzero entries, so each coordinate of $\hat{\mathbf{z}}$ mixes several factors (**one-to-many**). The matrix $\mathbf{A}$ is constructed via its singular value decomposition $\mathbf{A} = \mathbf{U} \text{diag}(s_1, \dots, s_d) \mathbf{V}^\top$ with $\mathbf{U} \in \mathbb{R}^{m \times m}$, $\mathbf{V} \in \mathbb{R}^{d \times d}$

orthogonal, and $s_1 \geq \dots \geq s_d > 0$. The condition number $\kappa(\mathbf{A}) = s_1/s_d$ controls the degree of entanglement. Since $\text{rank}(\mathbf{A}) = d$, the factor information is globally preserved; recovery requires r to unmix the linear superposition.

Example ($d = 2, m = 4$): every $\hat{z}_j$ is a distinct linear combination of z_1 and z_2 .

E8 — Distributed ($m > d$, nonlinear). Let $k \geq 2$ be an integer and set $m = k \cdot d$. There exist pairwise-disjoint index sets $S_1, \dots, S_d \subset \{1, \dots, m\}$ with $|S_i| = k$, $\bigsqcup_{i=1}^d S_i = \{1, \dots, m\}$, such that each ground-truth factor z_i is encoded *only* in the coordinates indexed by S_i (no cross-factor mixing):

$$\forall j \in S_i : \quad \hat{z}_j = h_j(z_{\pi(i)}),$$

where $\pi \in S_d$ is a permutation and each $h_j : \mathbb{R} \rightarrow \mathbb{R}$ is a scalar nonlinear function (not necessarily invertible individually). The factor is recoverable from its own subset via a decoder $f_i : \mathbb{R}^k \rightarrow \mathbb{R}$:

$$z_{\pi(i)} = f_i(\hat{z}_j : j \in S_i).$$

For $k = 2$, the canonical implementation uses $\hat{z}_{2i} = \sin(z_{\pi(i)})$, $\hat{z}_{2i+1} = \cos(z_{\pi(i)})$, with perfect reconstruction via $z_{\pi(i)} = \text{atan2}(\hat{z}_{2i}, \hat{z}_{2i+1})$. For $k > 2$, an interval-based encoding partitions the range of each factor into k bins; exactly one code per factor is active for each sample. Coordinate-wise metrics fail because the readout r must aggregate (*many-to-one*) across the k codes in each S_i ; no single $\hat{z}_j$ suffices to recover z_i .

Example ($d = 2, k = 2, m = 4$): $(\sin z_1, \cos z_1, \sin z_2, \cos z_2)$.

Control baselines. Every metric should return ≈ 0 ; any nonzero score is a false positive.

E9 — Random (independent of data). $\hat{\mathbf{z}} \sim \text{Normal}(0, I_m)$, independent of $\mathbf{x}$.

E10 — Random (independent of data). $\hat{\mathbf{z}} \sim \text{Uniform}([-1, 1]^m)$, independent of $\mathbf{x}$.

F METRICS

See § G for a more detailed description of each metric.

Metric	Short description
DCI	Measures <i>Disentanglement</i> (D), <i>Completeness</i> (C), and <i>Informativeness</i> (I) by assessing how well latent dimensions predict ground-truth factors using supervised regressors (Eastwood and Williams, 2018).
MCC	Evaluates alignment between learned and ground-truth latent variables via an optimal one-to-one matching that maximizes pairwise correlations (P=pearson, S=spearman, RDC=Randomized Dependence Coefficient (Lopez-Paz et al., 2013)).
R^2	Quantifies the proportion of variance in ground-truth factors explained by the learned representation through linear regression.
T-MEX	Assesses disentanglement by measuring how selectively latent variables respond to interventions on ground-truth factors (Yao et al., 2025).
MIG	Computes the gap between the top two mutual information scores between a factor and latent variables (Sreekar et al., 2021).
InfoMEC	Measures equivalence classes of representations by evaluating how much information about the ground-truth factors is preserved under invertible transformations (Hsu et al., 2023) (M=modularity, E=explicitness, C=compactness).

Table 2: Common evaluation metrics for causal representation learning with ground-truth factors.

Table 3: **No single metric satisfies all four properties.** Each cell shows whether a metric satisfies (✓) the property in $\geq 95\%$ of the tested (DGP, encoder) settings, partially satisfies ($\sim$) in 50 – 95% settings, or violates (✗) the corresponding property (so doesn’t satisfy the property most of the time except for exceptions). Superscripts reference the relevant subsection. See § H for MI-based metrics.

Metric	P1	P2	P3	P4
MCC-P	✗	✗	✗	✗
MCC-S	✗	✗	✗	✗
R^2	✓	$\sim$	$\sim$	✓
DCI-D	$\sim$	$\sim$	✗	$\sim$
MIG	✗	✗	✗	$\sim$
T-MEX	$\sim$	✗	✗	$\sim$

P1: ρ -invariance^{3.1} P2: d_{eff} -sensitivity^{3.2}
P3: OC-invariance^{3.3} P4: Uninformative-sensitivity^{3.4}

G EXPECTED METRICS’ BEHAVIOUR: THEORY AND DERIVATIONS

In this section, we derive the theoretically expected behavior of some metrics to highlight some of their limitations, even in simple settings.

G.1 MCC: CORRELATED LATENT FACTORS AND LINEAR ENTANGLEMENT

We consider three ground-truth latent variables

$$\mathbf{z} = (Z_1, Z_2, Z_3)^\top,$$

with the following second-order structure:

$$\mathbb{E}[Z_i] = 0 \quad \text{for all } i, \quad (1)$$

$$\sigma^2(Z_i) = 1 \quad \text{for all } i, \quad (2)$$

$$\text{Corr}(Z_2, Z_3) = \rho, \quad |\rho| < 1, \quad (3)$$

and Z_1 uncorrelated with (Z_2, Z_3) .

Consider a learned representation with linear mixing of the form:

$$\hat{Z}_1 = sZ_1, \quad (4)$$

$$\hat{Z}_2 = aZ_2 + bZ_3, \quad (5)$$

$$\hat{Z}_3 = cZ_2 + dZ_3, \quad (6)$$

with $s, a, b, c, d \neq 0$. Our goal is to compute the Mean Correlation Coefficient (MCC) between $\mathbf{Z}$ and $\hat{\mathbf{Z}} := (\hat{Z}_1, \hat{Z}_2, \hat{Z}_3)$ as a function of the latent correlation ρ .

We start by computing the covariances between the true latents and the learned coordinates. For Z_1 we immediately have

$$\text{Cov}(Z_1, \hat{Z}_1) = s\sigma^2(Z_1) = s, \quad (7)$$

$$\text{Cov}(Z_1, \hat{Z}_2) = 0, \quad (8)$$

$$\text{Cov}(Z_1, \hat{Z}_3) = 0, \quad (9)$$

since Z_1 is uncorrelated with Z_2 and Z_3 , and $\sigma^2(Z_1) = 1$.

For Z_2 and $\hat{Z}_2 = aZ_2 + bZ_3$,

$$\text{Cov}(Z_2, \hat{Z}_2) = \text{Cov}(Z_2, aZ_2 + bZ_3) \quad (10)$$

$$= a \text{Cov}(Z_2, Z_2) + b \text{Cov}(Z_2, Z_3) \quad (11)$$

$$= a\sigma^2(Z_2) + b\rho\sqrt{\sigma^2(Z_2)\sigma^2(Z_3)} \quad (12)$$

$$= a + b\rho, \quad (13)$$

using $\sigma^2(Z_2) = \sigma^2(Z_3) = 1$ and $\text{Cov}(Z_2, Z_3) = \rho$.

Similarly, the variance of $\hat{Z}_2$ is

$$\sigma^2(\hat{Z}_2) = \sigma^2(aZ_2 + bZ_3) \quad (14)$$

$$= a^2\sigma^2(Z_2) + b^2\sigma^2(Z_3) + 2ab\text{Cov}(Z_2, Z_3) \quad (15)$$

$$= a^2 + b^2 + 2ab\rho. \quad (16)$$

Therefore

$$\text{Corr}(Z_2, \hat{Z}_2) = \frac{\text{Cov}(Z_2, \hat{Z}_2)}{\sqrt{\sigma^2(Z_2)\sigma^2(\hat{Z}_2)}} = \frac{a + b\rho}{\sqrt{a^2 + b^2 + 2ab\rho}}. \quad (17)$$

Analogously,

$$\text{Cov}(Z_3, \hat{Z}_2) = a\text{Cov}(Z_3, Z_2) + b\text{Cov}(Z_3, Z_3) = a\rho + b, \quad (18)$$

$$\text{Corr}(Z_3, \hat{Z}_2) = \frac{b + a\rho}{\sqrt{a^2 + b^2 + 2ab\rho}}. \quad (19)$$

Repeating the same computation for $\hat{Z}_3 = cZ_2 + dZ_3$ yields

$$\text{Corr}(Z_2, \hat{Z}_3) = \frac{c + d\rho}{\sqrt{c^2 + d^2 + 2cd\rho}}, \quad (20)$$

$$\text{Corr}(Z_3, \hat{Z}_3) = \frac{d + c\rho}{\sqrt{c^2 + d^2 + 2cd\rho}}. \quad (21)$$

Finally, for Z_1 we have

$$\text{Corr}(Z_1, \hat{Z}_1) = \text{sgn}(s), \quad \text{Corr}(Z_1, \hat{Z}_2) = \text{Corr}(Z_1, \hat{Z}_3) = 0. \quad (22)$$

G.1.1 Correlation matrix and MCC

Collecting the correlations into a matrix $C(\rho)$, we obtain

$$C(\rho) = \begin{pmatrix} \text{Corr}(Z_1, \hat{Z}_1) & \text{Corr}(Z_1, \hat{Z}_2) & \text{Corr}(Z_1, \hat{Z}_3) \\ \text{Corr}(Z_2, \hat{Z}_1) & \text{Corr}(Z_2, \hat{Z}_2) & \text{Corr}(Z_2, \hat{Z}_3) \\ \text{Corr}(Z_3, \hat{Z}_1) & \text{Corr}(Z_3, \hat{Z}_2) & \text{Corr}(Z_3, \hat{Z}_3) \end{pmatrix} \quad (23)$$

$$= \begin{pmatrix} \pm 1 & 0 & 0 \\ 0 & r_{22}(\rho) & r_{23}(\rho) \\ 0 & r_{32}(\rho) & r_{33}(\rho) \end{pmatrix}, \quad (24)$$

where $r_{22}, r_{23}, r_{32}, r_{33}$ are given by (17)–(21).

The Mean Correlation Coefficient (MCC) between $\mathbf{z}$ and $\hat{\mathbf{z}}$ for the 3×3 case is defined as

$$\text{MCC}(\rho) = \frac{1}{3} \max_{\pi \in S_3} \sum_{i=1}^3 |\text{Corr}(Z_i, \hat{Z}_{\pi(i)})|, \quad (25)$$

where S_3 is the set of all permutations of $\{1, 2, 3\}$.

Since $\text{Corr}(Z_1, \hat{Z}_1) = \pm 1$ and $\text{Corr}(Z_1, \hat{Z}_2) = \text{Corr}(Z_1, \hat{Z}_3) = 0$, the optimal permutation always pairs Z_1 with $\hat{Z}_1$, contributing 1 to the sum. The remaining degrees of freedom are in how we pair (Z_2, Z_3) with $(\hat{Z}_2, \hat{Z}_3)$.

There are two relevant pairings:

1. “Diagonal” pairing: $Z_2 \leftrightarrow \hat{Z}_2$ and $Z_3 \leftrightarrow \hat{Z}_3$, giving a sum

$$S_{\text{diag}}(\rho) = |r_{22}(\rho)| + |r_{33}(\rho)|.$$

2. “Swapped” pairing: $Z_2 \leftrightarrow \hat{Z}_3$ and $Z_3 \leftrightarrow \hat{Z}_2$, giving

$$S_{\text{swap}}(\rho) = |r_{23}(\rho)| + |r_{32}(\rho)|.$$

Thus

$$\text{MCC}(\rho) = \frac{1}{3} \left(1 + \max\{S_{\text{diag}}(\rho), S_{\text{swap}}(\rho)\} \right). \quad (26)$$

The dependence of MCC on the latent correlation ρ is therefore entirely through the functions $r_{22}(\rho), \dots, r_{33}(\rho)$.

Effect of the sign of ρ in a symmetric example. To see how the sign of ρ affects $\text{MCC}(\rho)$, consider a symmetric mixing:

$$\hat{Z}_2 = Z_2 + \varepsilon Z_3, \quad (27)$$

$$\hat{Z}_3 = \varepsilon Z_2 + Z_3, \quad (28)$$

with $0 < \varepsilon < 1$. In this case

$$a = d = 1, \quad b = c = \varepsilon,$$

and substituting into (17)–(21) yields

$$r_{22}(\rho) = \frac{1 + \varepsilon\rho}{\sqrt{1 + \varepsilon^2 + 2\varepsilon\rho}}, \quad (29)$$

$$r_{33}(\rho) = \frac{1 + \varepsilon\rho}{\sqrt{1 + \varepsilon^2 + 2\varepsilon\rho}} = r_{22}(\rho), \quad (30)$$

and the off-diagonal entries are

$$r_{32}(\rho) = \frac{\varepsilon + \rho}{\sqrt{1 + \varepsilon^2 + 2\varepsilon\rho}}, \quad (31)$$

$$r_{23}(\rho) = \frac{\varepsilon + \rho}{\sqrt{\varepsilon^2 + 1 + 2\varepsilon\rho}} = r_{32}(\rho). \quad (32)$$

To verify: from (21) with $d = 1$, $c = \varepsilon$, the numerator is $d + c\rho = 1 + \varepsilon\rho$, matching r_{22} . From (19) with $a = 1$, $b = \varepsilon$, the numerator is $b + a\rho = \varepsilon + \rho$. From (20) with $c = \varepsilon$, $d = 1$, the numerator is $c + d\rho = \varepsilon + \rho$, matching r_{32} . All four denominators equal $\sqrt{1 + \varepsilon^2 + 2\varepsilon\rho}$.

The two pairings therefore give

$$S_{\text{diag}}(\rho) = |r_{22}(\rho)| + |r_{33}(\rho)| = 2|r_{22}(\rho)|, \quad (33)$$

$$S_{\text{swap}}(\rho) = |r_{23}(\rho)| + |r_{32}(\rho)| = 2|r_{32}(\rho)|. \quad (34)$$

These are *not* equal in general. We now show that the diagonal pairing always dominates. Consider the difference of the numerators:

$$(1 + \varepsilon\rho) - (\varepsilon + \rho) = (1 - \varepsilon)(1 - \rho) \geq 0, \quad (35)$$

$$(1 + \varepsilon\rho) + (\varepsilon + \rho) = (1 + \varepsilon)(1 + \rho) > 0, \quad (36)$$

for all $\rho \in (-1, 1)$ and $\varepsilon \in (0, 1)$. Since both share the same (positive) denominator, we have $|r_{22}(\rho)| \geq |r_{32}(\rho)|$ with equality only at the boundary $\rho = 1$ or $\varepsilon = 1$. Hence $S_{\text{diag}}(\rho) \geq S_{\text{swap}}(\rho)$, and

$$\text{MCC}(\rho) = \frac{1}{3} \left(1 + 2|r_{22}(\rho)| \right) = \frac{1}{3} \left(1 + \frac{2|1 + \varepsilon\rho|}{\sqrt{1 + \varepsilon^2 + 2\varepsilon\rho}} \right). \quad (37)$$

Since $1 + \varepsilon\rho \geq 1 - \varepsilon > 0$ for all $|\rho| < 1$, the absolute value is redundant and we may write

$$\text{MCC}(\rho) = \frac{1}{3} \left(1 + \frac{2(1 + \varepsilon\rho)}{\sqrt{1 + \varepsilon^2 + 2\varepsilon\rho}} \right). \quad (38)$$

Derivative of $r_{22}(\rho)$. We now compute the derivative of $r_{22}(\rho)$ with respect to ρ . Write $N(\rho) := 1 + \varepsilon\rho$ and $D(\rho) := 1 + \varepsilon^2 + 2\varepsilon\rho$, so that $r_{22} = N/\sqrt{D}$. Then

$$\frac{\partial r_{22}}{\partial \rho} = \frac{N' \sqrt{D} - N \cdot \frac{D'}{2\sqrt{D}}}{D} = \frac{N' D - N \cdot \frac{D'}{2}}{D^{3/2}}. \quad (39)$$

We have $N' = \varepsilon$ and $D' = 2\varepsilon$, so the numerator is

$$\begin{aligned} N' D - N \cdot \frac{D'}{2} &= \varepsilon(1 + \varepsilon^2 + 2\varepsilon\rho) - (1 + \varepsilon\rho) \cdot \varepsilon \\ &= \varepsilon[(1 + \varepsilon^2 + 2\varepsilon\rho) - (1 + \varepsilon\rho)] \\ &= \varepsilon(\varepsilon^2 + \varepsilon\rho) = \varepsilon^2(\varepsilon + \rho). \end{aligned} \quad (40)$$

Therefore

$$\frac{\partial r_{22}}{\partial \rho} = \frac{\varepsilon^2(\varepsilon + \rho)}{(1 + \varepsilon^2 + 2\varepsilon\rho)^{3/2}}. \quad (41)$$

The denominator in (41) is strictly positive for all $\rho \in (-1, 1)$ and $0 < \varepsilon < 1$, since

$$1 + \varepsilon^2 + 2\varepsilon\rho \geq 1 + \varepsilon^2 - 2\varepsilon = (1 - \varepsilon)^2 > 0.$$

Non-monotonicity of $r_{22}(\rho)$. The numerator in (41) changes sign at $\rho^* = -\varepsilon$:

$$\frac{\partial r_{22}}{\partial \rho} \begin{cases} < 0 & \text{if } \rho < -\varepsilon, \\ = 0 & \text{if } \rho = -\varepsilon, \\ > 0 & \text{if } \rho > -\varepsilon. \end{cases} \quad (42)$$

Thus $r_{22}(\rho)$ is *not* monotonically increasing on $(-1, 1)$. It attains its minimum at $\rho^* = -\varepsilon$, where

$$r_{22}(-\varepsilon) = \frac{1 - \varepsilon^2}{\sqrt{1 + \varepsilon^2 - 2\varepsilon^2}} = \frac{1 - \varepsilon^2}{\sqrt{1 - \varepsilon^2}} = \sqrt{1 - \varepsilon^2}. \quad (43)$$

Monotonicity of MCC. Since $1 + \varepsilon\rho > 0$ on $(-1, 1)$, we have $|r_{22}| = r_{22}$, and by (38) the MCC inherits the same monotonicity structure: decreasing on $(-1, -\varepsilon)$ and increasing on $(-\varepsilon, 1)$, with minimum

$$\text{MCC}(-\varepsilon) = \frac{1}{3}(1 + 2\sqrt{1 - \varepsilon^2}). \quad (44)$$

At the boundary values:

$$\lim_{\rho \rightarrow 1} r_{22}(\rho) = \frac{1 + \varepsilon}{\sqrt{(1 + \varepsilon)^2}} = 1, \quad (45)$$

$$\lim_{\rho \rightarrow -1^+} r_{22}(\rho) = \frac{1 - \varepsilon}{\sqrt{(1 - \varepsilon)^2}} = 1. \quad (46)$$

Hence $r_{22}(\rho) \rightarrow 1$ at *both* extremes $\rho \rightarrow \pm 1$, and $\text{MCC}(\rho) \rightarrow 1$ in both limits. The minimum of MCC is in the interior, at $\rho = -\varepsilon$.

Implication. Even though $\hat{Z}_2$ and $\hat{Z}_3$ remain linearly entangled mixtures of Z_2 and Z_3 for all $\varepsilon > 0$, the MCC score varies with the correlation ρ between the ground-truth factors, despite the underlying entanglement structure of the learned representation being unchanged.

For the practically relevant regime $\rho > 0$, positive correlations inflate MCC monotonically: the MCC is strictly increasing on $(0, 1)$ since $0 > -\varepsilon = \rho^*$. For negative correlations, the MCC first decreases (reaching its minimum at $\rho = -\varepsilon$) and then increases again toward 1 as $\rho \rightarrow -1$.

The non-trivial dependence on ρ —including the fact that the minimum is in the interior and that the MCC approaches 1 at *both* boundary values $\rho \rightarrow \pm 1$ —demonstrates that MCC conflates representation quality with the covariance structure of the ground-truth factors.

Correlation (ρ) vs Entanglement ($\kappa \in \{1, 2, \dots, 50\}$)
(D2 + E3, $d=2$, $n=100$)

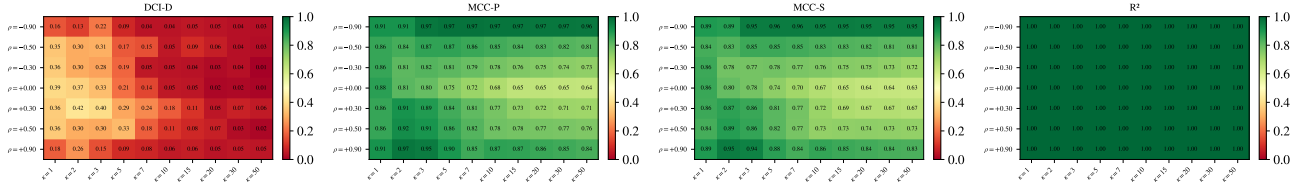


Figure 6: MCC overestimates identifiability as ρ increases.

To illustrate this further, consider an even simpler (degenerate) example:

$$a = b = c = d = 1, \quad (47)$$

so that

$$\hat{Z}_2 = Z_2 + Z_3, \quad \hat{Z}_3 = Z_2 + Z_3. \quad (48)$$

In other words, $\hat{Z}_2$ and $\hat{Z}_3$ are identical, fully redundant, and both are symmetric mixtures of Z_2 and Z_3 .

Using the same covariance structure as before, with

$$\sigma^2(Z_2) = \sigma^2(Z_3) = 1, \quad \text{Cov}(Z_2, Z_3) = \rho,$$

we compute

$$\begin{aligned} \text{Cov}(Z_2, \hat{Z}_2) &= \text{Cov}(Z_2, Z_2 + Z_3) = \sigma^2(Z_2) + \text{Cov}(Z_2, Z_3) = 1 + \rho, \\ \sigma^2(\hat{Z}_2) &= \sigma^2(Z_2 + Z_3) = \sigma^2(Z_2) + \sigma^2(Z_3) + 2 \text{Cov}(Z_2, Z_3) = 2 + 2\rho. \end{aligned}$$

Hence, for $\rho > -1$,

$$\text{Corr}(Z_2, \hat{Z}_2) = \frac{1 + \rho}{\sqrt{2 + 2\rho}} = \sqrt{\frac{1 + \rho}{2}} =: r(\rho). \quad (49)$$

Verification via the general formula. Setting $\varepsilon = 1$ in (29) gives $r_{22} = (1 + \rho)/\sqrt{2 + 2\rho}$, matching (49).

By symmetry we also have

$$\text{Corr}(Z_3, \hat{Z}_2) = \text{Corr}(Z_2, \hat{Z}_2) = \text{Corr}(Z_2, \hat{Z}_3) = \text{Corr}(Z_3, \hat{Z}_3) = r(\rho),$$

and $\text{Corr}(Z_1, \hat{Z}_1) = \pm 1$, $\text{Corr}(Z_1, \hat{Z}_2) = \text{Corr}(Z_1, \hat{Z}_3) = 0$ as before.

In this case, the two candidate pairings (diagonal and swapped) give the same sum, and the MCC simplifies to

$$\text{MCC}(\rho) = \frac{1}{3}(1 + 2r(\rho)) = \frac{1}{3}\left(1 + 2\sqrt{\frac{1 + \rho}{2}}\right), \quad -1 < \rho \leq 1. \quad (50)$$

Note that this degenerate case corresponds to $\varepsilon = 1$, where the minimum of r_{22} from (42) occurs at $\rho^* = -1$ (the boundary), consistent with $r(\rho)$ being monotonically increasing on $(-1, 1)$.

This expression makes two important properties explicit.

(i) Asymmetry $\text{MCC}(\rho) \neq \text{MCC}(-\rho)$. From (50) we obtain

$$r(\rho) = \sqrt{\frac{1 + \rho}{2}}, \quad -1 < \rho \leq 1. \quad (51)$$

Substituting $-\rho$ in place of ρ yields

$$\text{MCC}(-\rho) = \frac{1}{3}\left(1 + 2\sqrt{\frac{1 - \rho}{2}}\right). \quad (52)$$

Hence, for any $\rho \in (0, 1)$,

$$\text{MCC}(\rho) = \frac{1}{3} \left(1 + 2\sqrt{\frac{1+\rho}{2}} \right) \neq \frac{1}{3} \left(1 + 2\sqrt{\frac{1-\rho}{2}} \right) = \text{MCC}(-\rho).$$

Even though the entanglement structure for (Z_2, Z_3) is symmetric under the sign flip $\rho \mapsto -\rho$, MCC values differ for positive and negative correlations.

(ii) Faithfulness issues at extreme correlations. The same formula reveals a qualitative difference between the limits $\rho \rightarrow 1$ and $\rho \rightarrow -1$:

$$\lim_{\rho \rightarrow 1} r(\rho) = \sqrt{\frac{1+1}{2}} = 1, \quad (53)$$

$$\lim_{\rho \rightarrow -1^+} r(\rho) = \sqrt{\frac{1+(-1)}{2}} = 0. \quad (54)$$

When $\rho \rightarrow 1$, Eq. (50) yields $\text{MCC}(\rho) \rightarrow (1+2)/3 = 1$. In contrast, when $\rho \rightarrow -1$, $\text{MCC}(\rho) \rightarrow (1+0)/3 \approx 0.33$.

Contrast with the $\varepsilon < 1$ case. In the general symmetric mixing with $\varepsilon < 1$, eqs. (45)–(46) show that $r_{22}(\rho) \rightarrow 1$ at *both* $\rho \rightarrow +1$ and $\rho \rightarrow -1$, so $\text{MCC}(\rho) \rightarrow 1$ in both limits. The faithfulness collapse ($\text{MCC} \rightarrow 0.33$) at $\rho \rightarrow -1$ is specific to the degenerate case $\varepsilon = 1$, where $\hat{Z}_2 = \hat{Z}_3 = Z_2 + Z_3 \approx 0$ when $Z_3 \approx -Z_2$. For $\varepsilon < 1$, the representations $\hat{Z}_2 \neq \hat{Z}_3$ remain distinct and their correlations with the true factors recover to 1 as $\rho \rightarrow -1$.

These examples show that MCC depends nontrivially on the covariance structure of the ground-truth factors, independently of the underlying entanglement of the learned representation. In particular, for a fixed mixing matrix, MCC can be artificially inflated or deflated by the latent correlation ρ .

G.2 EXPECTED BEHAVIOUR OF DCI

We derive properties of DCI that explain the metric’s behaviour in the main-text experiments. We first recall the construction, then state four results organised by failure mode.

Construction. A supervised probe is trained to predict each ground-truth factor z_j from the learned representation $\hat{\mathbf{z}}$, yielding a nonnegative importance matrix $\mathcal{R} \in \mathbb{R}_{\geq 0}^{m \times d}$, where $R_{i,j}$ quantifies the contribution of learned feature $\hat{z}_i$ in predicting z_j . Row $\mathcal{R}_{i,:}$ summarises which factors feature i encodes; column $\mathcal{R}_{:,j}$ summarises which features encode factor j . The DCI scores are computed from $\mathcal{R}$ alone:

Disentanglement. Convert each row to a distribution $p_{j|i} = R_{i,j} / \sum_k R_{i,k}$ and measure concentration:

$$D_i = 1 - \frac{H(p_{\cdot|i})}{\log d}, \quad D_{\text{DCI}} = \sum_{i=1}^m w_i D_i, \quad w_i = \frac{\sum_j R_{i,j}}{\sum_{i',j'} R_{i',j'}}.$$

Completeness. Convert each column to a distribution $\tilde{p}_{i|j} = R_{i,j} / \sum_k R_{k,j}$ and measure concentration:

$$C_j = 1 - \frac{H(\tilde{p}_{\cdot|j})}{\log m}, \quad C_{\text{DCI}} = \sum_{j=1}^d v_j C_j, \quad v_j = \frac{\sum_i R_{i,j}}{\sum_{i',j'} R_{i',j'}}.$$

Informativeness. $I_{\text{DCI}} = \frac{1}{d} \sum_j (1 - L_j)$, where L_j is the normalised prediction loss (e.g., $1 - R_j^2$) of the probe for factor j .

The weights w_i and v_j are proportional to total importance: features or factors with negligible importance contribute negligibly to the global scores. This weighting is the source of the first failure mode.

Proposition 2 (Dropped factors are invisible to DCI). *Under $\mathbf{D}_\perp + \mathbf{E4}$ with $|S| = m < d$ perfectly identified factors, as $n \rightarrow \infty$: $D_{\text{DCI}} \rightarrow 1$ and $C_{\text{DCI}} \rightarrow 1$.*

Proof. For retained factors ($j \in S$), the encoder is elementwise, so the probe importance concentrates on a single coordinate: $p_{j|i}$ and $\tilde{p}_{i|j}$ are one-hot for the matched pair, giving $D_i = 1$ and $C_j = 1$.

For discarded factors ($j \notin S$), no learned feature predicts them: $R_{i,j} \approx 0$ for all i . Their weight $v_j \propto \sum_i R_{i,j} \approx 0$ vanishes from C_{DCI} . Likewise, the rows corresponding to codes that encode only retained factors carry all the weight in D_{DCI} .

DCI does not verify that *all* factors are represented; factors that are never encoded produce zero importance, vanish from the weighted averages, and do not penalise the score. $\square$

Implication: This explains the DCI-D false positive in Fig. 3 (left, $\mathbf{D}_1$): as factors are dropped, D or C do not penalise omission (only I_{DCI} would drop).

Proposition 3 (Functional dependence decreases D under a perfect encoder). *Under $\mathbf{D}_f$ with $z_2 = f(z_1)$ (deterministic) and a perfect elementwise encoder $\mathbf{EI}$ ($\hat{z}_j = a_j z_j$), a nonlinear probe (e.g., gradient boosted trees) yields $D_{\text{DCI}} < 1$.*

Proof. Since $z_2 = f(z_1)$ exactly, $\hat{z}_1 = a_1 z_1$ perfectly determines z_2 via f , so the nonlinear probe assigns $R_{1,2} > 0$ in addition to $R_{1,1} > 0$. Symmetrically, when f is invertible (e.g., $f(z) = z^3$), $\hat{z}_2 = a_2 f(z_1)$ determines z_1 via f^{-1} , so $R_{2,1} > 0$ and $D_2 < 1$. The remaining $d - 2$ codes are independent and achieve $D_i = 1$, but the deflated D_1 and D_2 pull down D_{DCI} through their nonzero weights $w_1, w_2 > 0$. $\square$

With a linear probe (e.g., Lasso), the result can differ. For $z_1 \sim \mathcal{N}(0, 1)$ and $z_2 = z_1^3$, the population normal equations yield zero cross-coefficients— z_1 and z_1^3 are linearly orthogonal under Gaussian moments—so $\mathcal{R}$ is diagonal and $D_{\text{DCI}} = 1$. The deflation under $\mathbf{D}_f$ is therefore *probe-dependent*: it arises only when the probe is expressive enough to detect the functional relationship f .

Implication. This explains the DCI-D dip in the $\mathbf{D}_f$ panels of Fig. 3 (right): even at $m = d$ (dashed line), DCI-D is below 1.0 because the functional constraint between z_1 and z_2 spreads importance across the corresponding codes.

G.3 MCC FALSE-POSITIVE RATE UNDER NULL ENCODERS

We derive the expected behaviour of MCC-P when the learned representation is independent of the ground-truth factors, explaining the inflation observed in Fig. 5.

Setup. Let $\mathbf{z} \in \mathbb{R}^d$ and $\hat{\mathbf{z}} \in \mathbb{R}^m$ be independent random vectors (null encoder), and let $(z^{(1)}, \hat{z}^{(1)}), \dots, (z^{(n)}, \hat{z}^{(n)})$ be n i.i.d. paired samples. The sample Pearson correlation between $\hat{z}_i$ and z_j is

$$\hat{\rho}_{ij} = \frac{\sum_{t=1}^n (\hat{z}_i^{(t)} - \bar{\hat{z}}_i)(z_j^{(t)} - \bar{z}_j)}{\sqrt{\sum_t (\hat{z}_i^{(t)} - \bar{\hat{z}}_i)^2} \sqrt{\sum_t (z_j^{(t)} - \bar{z}_j)^2}}.$$

Since $\hat{\mathbf{z}} \perp \mathbf{z}$, the true correlation is $\rho_{ij} = 0$ for all i, j .

Distribution of sample correlations under the null. For bivariate normal data with $\rho = 0$, the sample correlation satisfies

$$\frac{\hat{\rho} \sqrt{n-2}}{\sqrt{1-\hat{\rho}^2}} \sim t_{n-2} \quad (55)$$

exactly (Fisher, 1922). For non-Gaussian data, the exact t -distribution does not hold, but the asymptotic result $\sqrt{n} \hat{\rho} \xrightarrow{d} \mathcal{N}(0, 1)$ follows from the Central Limit Theorem (CLT) (Hoeffding, 1992). In either case, for large n ,

$$\hat{\rho}_{ij} \overset{\text{approx}}{\sim} \mathcal{N}\left(0, \frac{1}{n}\right). \quad (56)$$

Maximum absolute correlation. MCC-P computes the $m \times d$ matrix of absolute sample correlations $|\hat{\rho}_{ij}|$ and applies Hungarian matching to find the optimal one-to-one assignment.

Consider a single column j . The entries $\{|\hat{\rho}_{ij}|\}_{i=1}^m$ are approximately half-normal with scale $1/\sqrt{n}$. (They are not exactly independent—they share the $z_j^{(t)}$ samples—but the dependence is weak under the null since the $\hat{z}_i$ are independent across rows; see [Cai and Jiang \(2011\)](#) for rigorous derivations, especially Theorem 2). The maximum of m such entries satisfies, by standard extreme value theory for Gaussian maxima,

$$\mathbb{E} \left[\max_{i=1}^m |\hat{\rho}_{ij}| \right] \approx \sqrt{\frac{2 \log m}{n}}. \quad (57)$$

Here, the leading term carries a multiplicative factor, which is a small correction for practical values of m . Please refer to [Cai and Jiang \(2011\)](#) for more details.

MCC-P under the null. To lower-bound the Hungarian matching, consider a greedy assignment: assign column 1 its best row, remove that row, assign column 2 its best among the remaining $m - 1$ rows, and so on. This produces a valid one-to-one assignment, and column j selects from $m - j + 1$ remaining candidates. When $m \gg d$, every column still has $\approx m$ candidates, and the greedy score is close to the average column-wise maximum. Since the Hungarian matching is optimal over all one-to-one assignments, it scores at least as high as the greedy, giving

$$\mathbb{E}[\text{MCC-P}] \gtrsim \frac{1}{d} \sum_{j=1}^d \sqrt{\frac{2 \log(m - j + 1)}{n}} \approx \sqrt{\frac{2 \log m}{n}} \quad \text{when } m \gg d. \quad (58)$$

This bound is non-negligible whenever $\log m/n$ is not small. In practice this inflation is substantial even at moderate ratios:

m/n	$\sqrt{2 \log m/n}$ (bound)	MCC-P (observed, $m/d = 1$)
0.1	0.21	~ 0.3
0.2	0.30	~ 0.5
0.5	0.48	~ 0.83
1.0	0.68	~ 0.95

The bound captures the correct scaling: it explains why m/n governs the false-positive rate. But, it underestimates the magnitude at practical sample sizes for two reasons: (i) the extreme value approximation is loose at small m ; and (ii) at small n , the exact null distribution of $\hat{\rho}$ follows a scaled t_{n-2} (Eq. (55)), which has heavier tails than the Gaussian, pushing the maximum correlation above the asymptotic prediction.

Extension to MCC-S. Spearman correlation is the Pearson correlation applied to ranks. Under independence, the sample Spearman correlation also satisfies $\hat{\rho}_{ij}^S \approx N(0, 1/n)$ ([Hotelling and Pabst, 1936](#)), so the extreme value argument above applies verbatim: the $\sqrt{2 \log m/n}$ floor governs both MCC-P and MCC-S.

Why m/n governs and m/d does not. The bound (57) depends on m (candidates per column) and n (sample size), but not on d (number of columns). Adding more ground-truth factors adds more columns to the matching problem but does not change the distribution of each column’s maximum. The MCC averaging divides by d , but since each column contributes approximately the same expected maximum, the average is $\approx \sqrt{2 \log m/n}$ regardless of d . This is consistent with the empirical observation in Fig. 5: reading along rows (fixed m/d , varying m/n), scores increase; reading along columns (fixed m/n , varying m/d), scores are approximately constant.

Comparison with R^2 and DCI-D. R^2 uses cross-validated nonlinear regression, which does not exploit the maximum over candidates: it predicts each factor independently and averages the explained variance. Under the null, the cross-validated $R_j^2 \approx 0$ for each factor (overfitting is penalised by the held-out evaluation), so $R^2 \approx 0$ regardless of m/n .

DCI-D trains a Lasso probe for each factor. Under the null, the ℓ_1 penalty shrinks most coefficients to zero, but a few features can be spuriously selected—particularly when m is large relative to n ([Bühlmann and Van De Geer, 2011](#)). The resulting importance matrix has most entries near zero; the moderate inflation at high m/n and low m/d in Fig. 5 is consistent with a small number of spuriously selected features spreading enough importance mass to inflate the disentanglement score.

H EXPERIMENTS

Plot legends label DGP types by index: $\mathbf{D1} = \mathbf{D}_{\perp}$ (independent), $\mathbf{D2} = \mathbf{D}_{\rho}$ (correlated), $\mathbf{D3} = \mathbf{D}_f$ (redundant), $\mathbf{D4} = \mathbf{D}_F$ (synergistic).

H.1 SANITY CHECKS

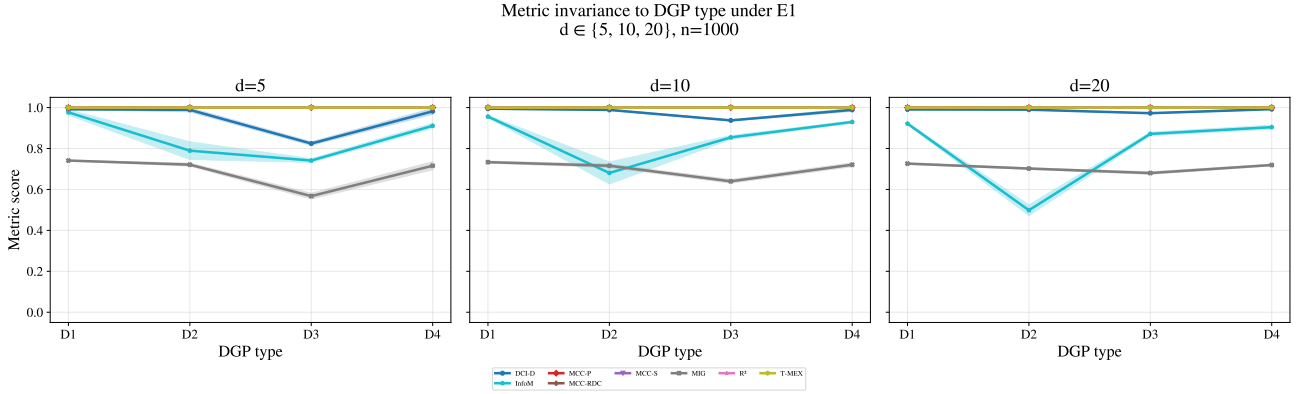


Figure 7: **Invariance to the DGP.** We expect a metric to be invariably equal to 1 for all DGPs under the ideal encoder E1. DCI-D is not invariant: the case $\mathbf{D}_f$ influences the metric. More strikingly, MIG and InfoM both report < 1 even in the ideal $\mathbf{D}_{\perp}$ –E1 case.

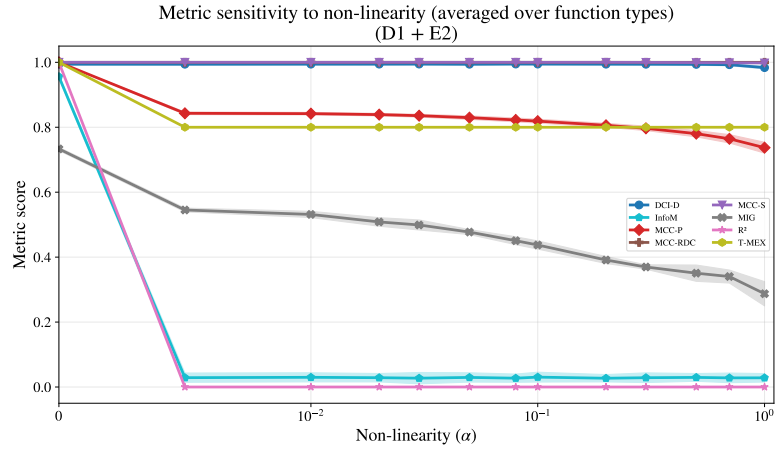


Figure 8: **Stability against nonlinear encoders.** DCI-D and MCC-S are stable against increasing non-linearity strength, hence reliable for encoder E2. On the other hand, MI-based metrics (MIG, InfoM) exhibit higher sensitivity to non-linearity. Metrics that explicitly assume linear transformations (MCC-P, R^2) are also sensitive, as expected. The list of considered nonlinear functions can be found in § E.2.

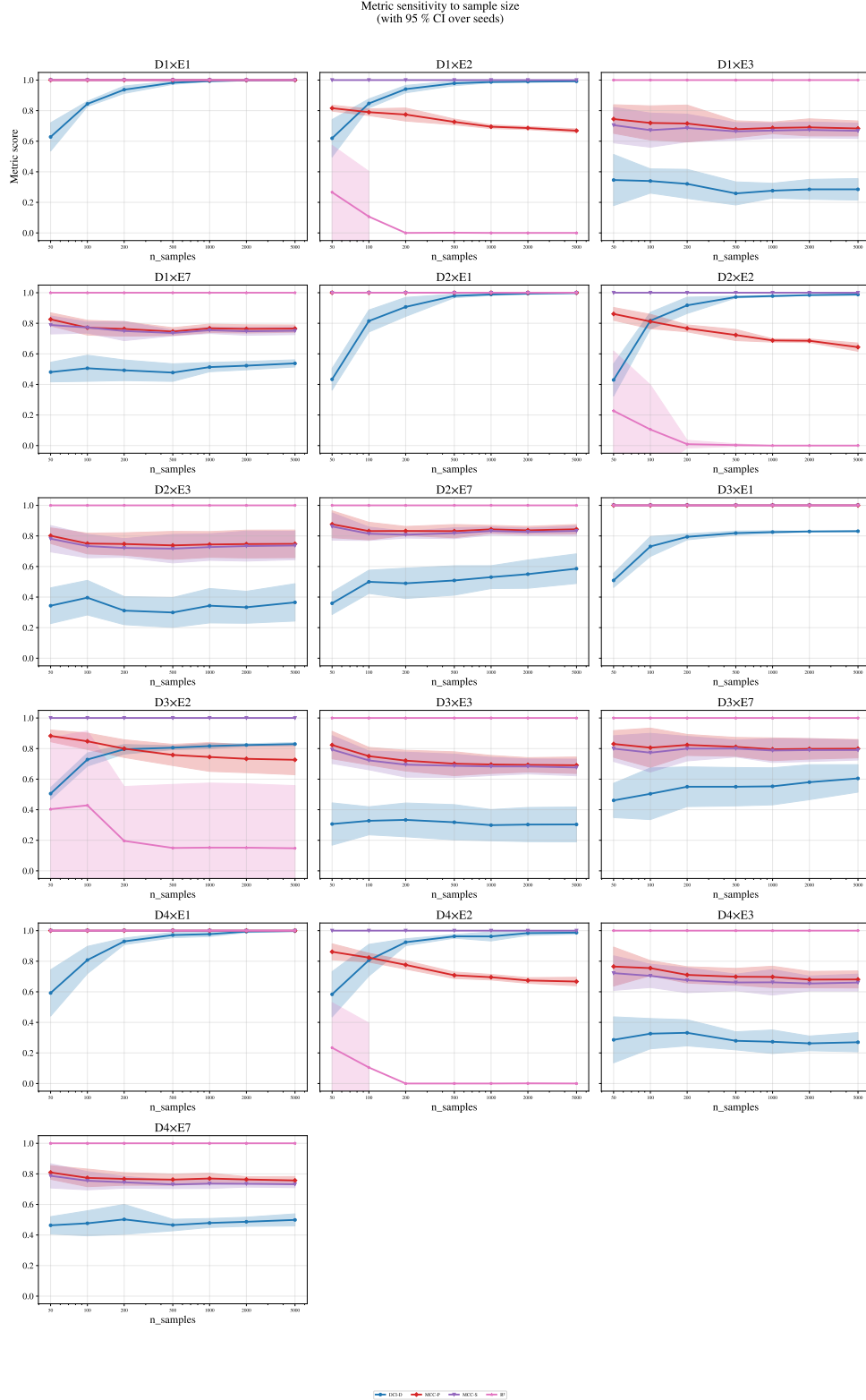


Figure 9: **Sample sensitivity** of four main metrics (DCI-D, MCC-P, MCC-S, R^2) across all DGP and encoder combinations ($n \in 50, \dots, 5000$, $d = 5$). MCC-P and MCC-S are sample-efficient: they stabilise by $n = 100$ in most settings. DCI-D requires $n \gtrsim 500$ to converge, with wide confidence intervals at $n < 200$, particularly under entangled encoders (E2, E3). The overcomplete encoder E7 ($m = 2d = 10$) inflates DCI-D variance at low n because twice as many codes must be estimated from the same number of samples.

Metric sensitivity to sample size
(with 95 % CI over seeds)

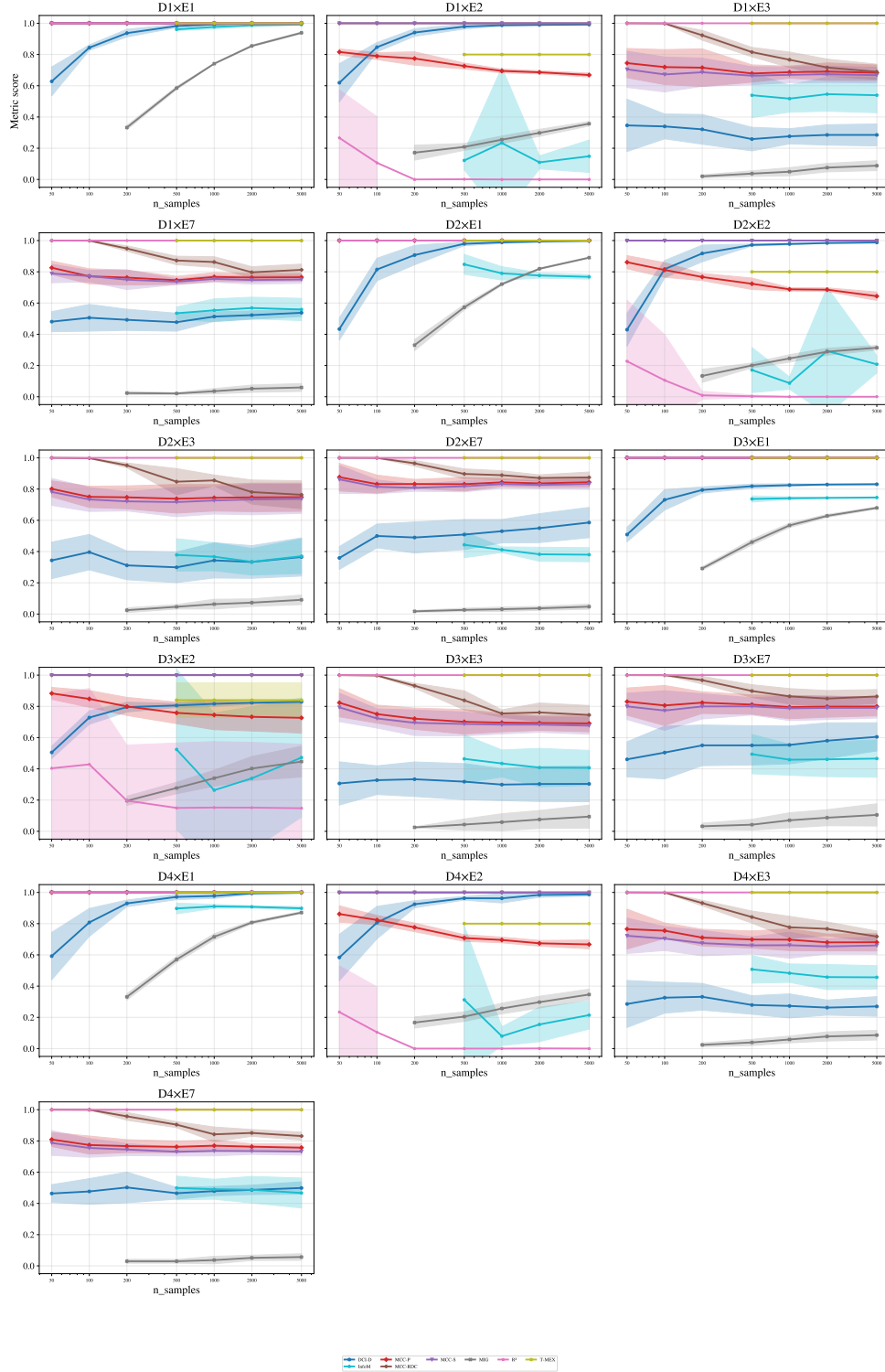


Figure 10: **Sample sensitivity.** Full metric suite (adding InfoM, MIG, MCC-RDC, T-MEX to the main four). T-MEX, InfoM and MIG produce NaN at $n = 50$ across all DGP and encoder combinations (visible as missing points), making them unusable below $n \approx 100$, and sometimes higher. MCC-RDC converges slowly under nonlinear encoders (**E2**), lagging behind MCC-P and MCC-S until $n \gtrsim 1000$. We observe a particularly high variance of the R^2 metric under **E2** (std ≈ 0.35 at $n = 50$), while correlation-based metrics remain stable (std < 0.01). Overall, metrics split into two reliability tiers: correlation-based measures (MCC-P, MCC-S) are robust across sample sizes, while predictor-based (R^2 , DCI-D) and information-theoretic (InfoM, MIG, T-MEX) metrics require $n \gtrsim 500$ for trustworthy estimates.

H.2 CORRELATION AMONG LATENT FACTORS

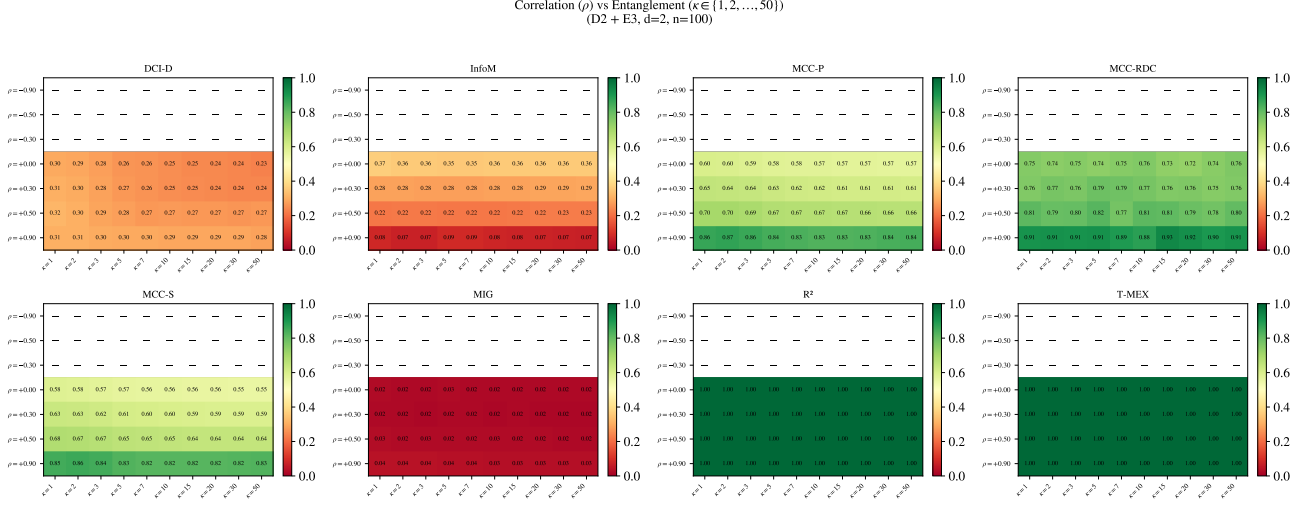


Figure 11: **Disentangling the effects of correlation ρ and entanglement κ ($d = 10$).** An ideal metric would vary only along columns (increasing κ , i.e. worse entanglement) and be constant along rows (changing ρ). All MCC variants and InfoM show row-wise gradients, confirming violation of Property 1. DCI-D is more stable but collapses to near-zero even under moderate entanglement (for all $\kappa > 1$).

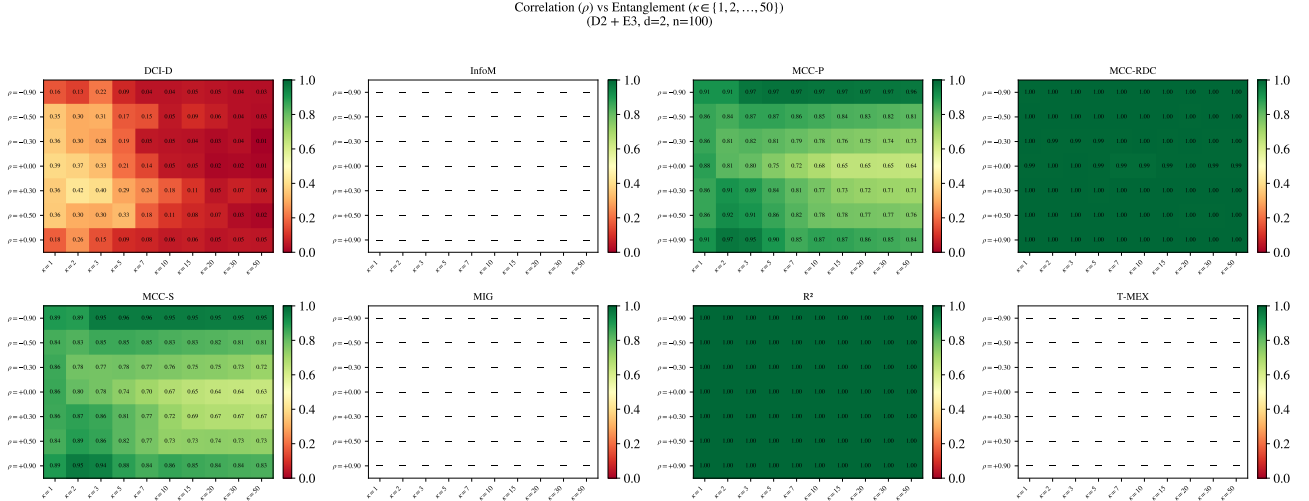


Figure 12: **Full metric suite: ρ - κ heatmaps at $d = 5$.** Qualitative conclusions carry over from $d = 10$ to $d = 5$, except for MCC-RDC which showcases a more stable behaviour. $n = 1000$.

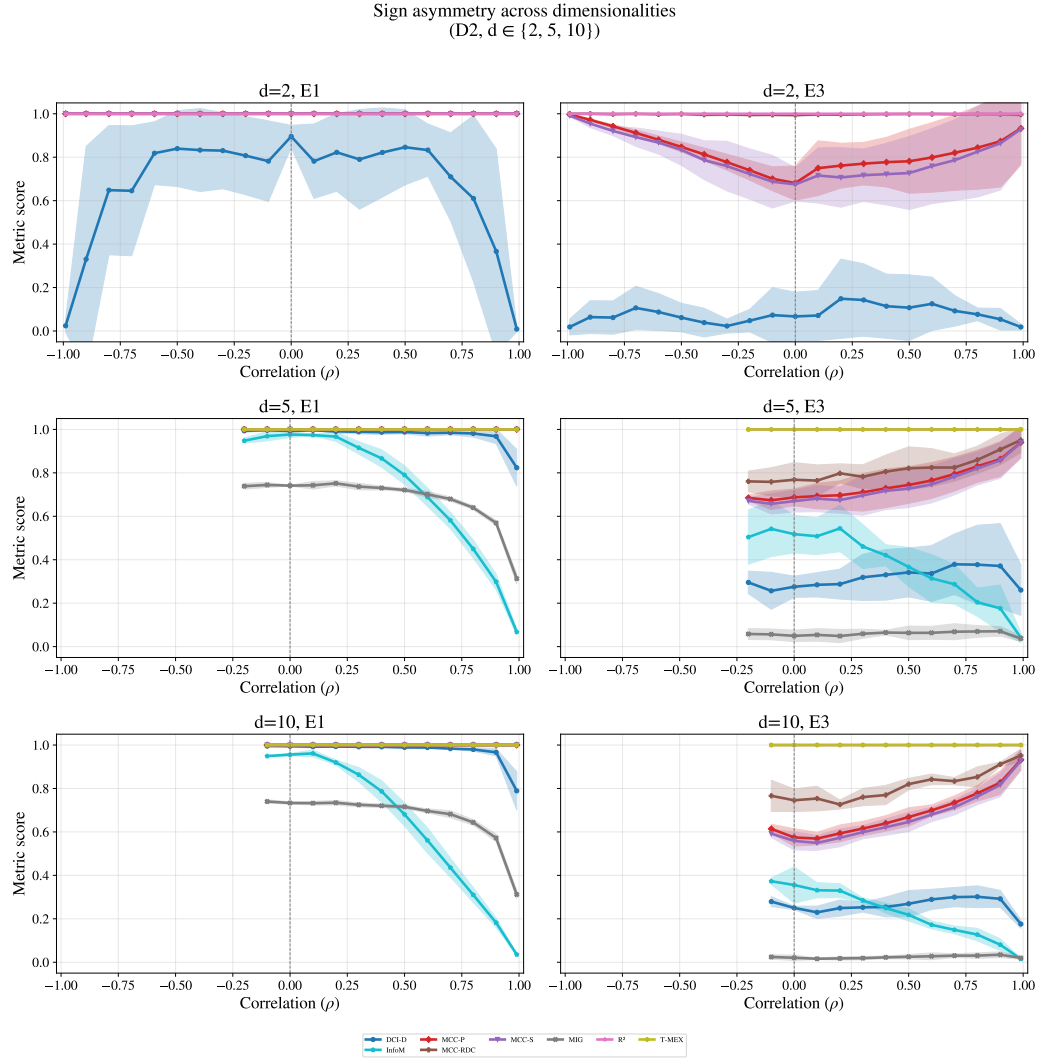


Figure 13: **Behaviour against correlation (D_ρ) and linear entanglement ($E3$), across dimensionalities, continuing Fig. 2.** The additional metrics reveal two distinct behaviours. MIG, InfoM, and DCI-D degrade with increasing $|\rho|$ under **E1**, indicating that these metrics conflate inter-factor correlation with non-identifiability. This tendency is also observed under **E3**, although linear entanglement is rightfully strongly penalized by them. On the other hand, under **E3**, all variants of MCC show an increase with inter-factor correlation, while we expected invariance: they conflate inter-factor correlation with disentanglement. MIG, InfoM, and T-MEX are absent at $d = 2$ because of numerical instabilities (NaN). R^2 and T-MEX remain flat near 1.0 with low variance under both **E1** and **E3** regardless of ρ , showing complete sign- and correlation-invariance, which also means that they are unable to flag linear correlation. Whether this is a desirable behaviour depends on the application.

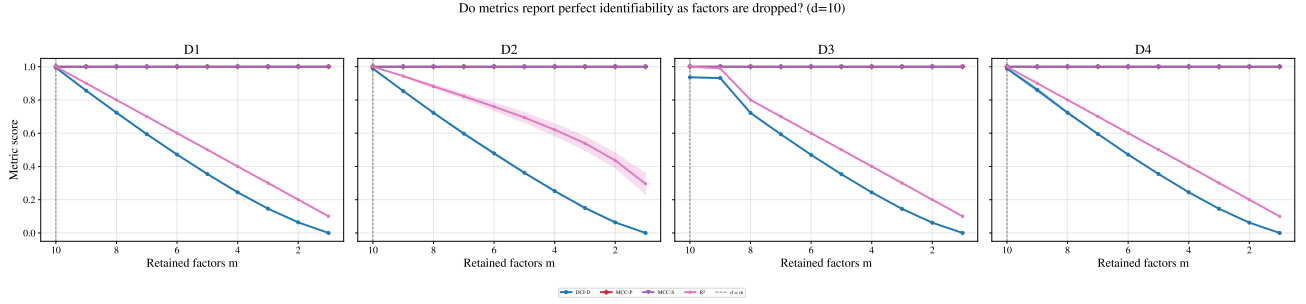


Figure 14: **Metric scores as a function of the number of retained factors across all DGP types.** Extension of Fig. 3 to $\mathbf{D}_\rho$ and $\mathbf{D}_F$. Under $\mathbf{D}_\rho$, R^2 exceeds m/d because the probe partially predicts dropped factors from correlated retained ones. Under $\mathbf{D}_F$ ($z_k = g(z_i, z_j)$, $d_{\text{eff}} = d-1$), metric behaviour is indistinguishable from $\mathbf{D}_\perp$ despite the first omission being lossless: no metric detects the multi-factor redundancy. MCC-P/S remain at 1.0 across all panels: dropping factors is not penalized by these metrics, as expected in theory (whether this is a desirable behaviour depends on the application). $d = 10$, $n = 1000$.

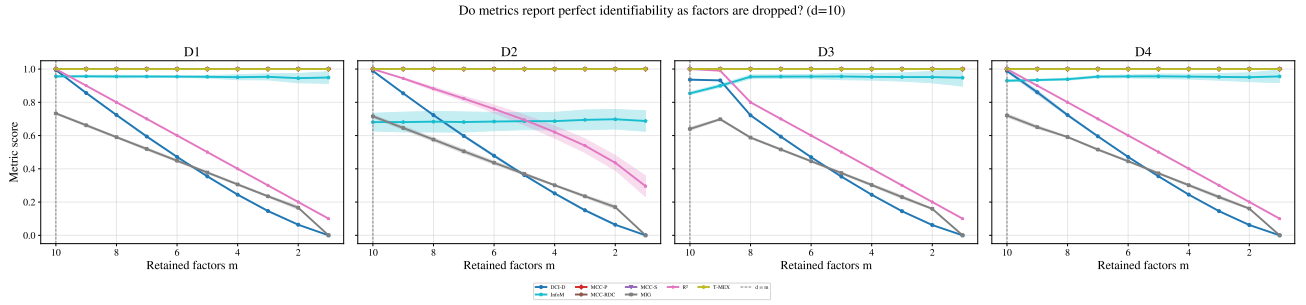


Figure 15: **Full metric suite for the dropped-factor experiment across all DGP types.** Extension of Fig. 14 with additional metrics. T-MEX, MCC-RDC and InfoM follow the behaviour of MCC-P/S, with InfoM sensitive to the correlation in $\mathbf{D}_\rho$. On the other hand, MIG behaves similarly as R^2 and DCI-D.

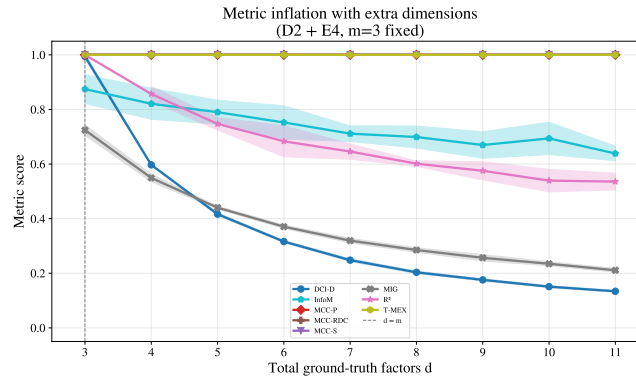


Figure 16: **Effect of inflating the number of ground-truth factors under $\mathbf{D}_\rho$.** The representation dimension is fixed at $m = 3$ while d increases by adding duplicated ground-truth factors. MCC-P/S/RDC remain constant as they match only m codes; R^2 and DCI-D decline because the probe must predict an increasing number of factors from the same m codes.

Overcomplete representations: E3 vs E5–E8
($d=20, n=1600$)

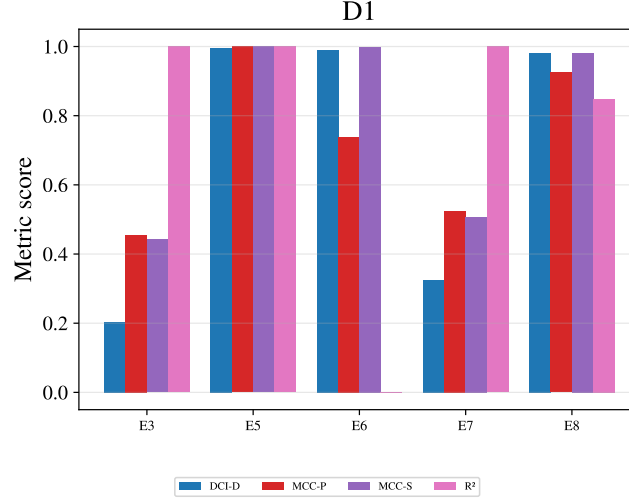


Figure 17: At moderate overcompleteness ($m/d = 2$), metrics distinguish entanglement from redundancy. Overcomplete disentangled encoders (E5, E6, E8) mostly score near 1.0 on DCI-D, MCC, and R^2 , whereas the entangled encoders (E3, E7) are correctly penalised by DCI-D and MCC. $m = 40, d = 20, n = 1600$. See Fig. 18 for all DGPs.

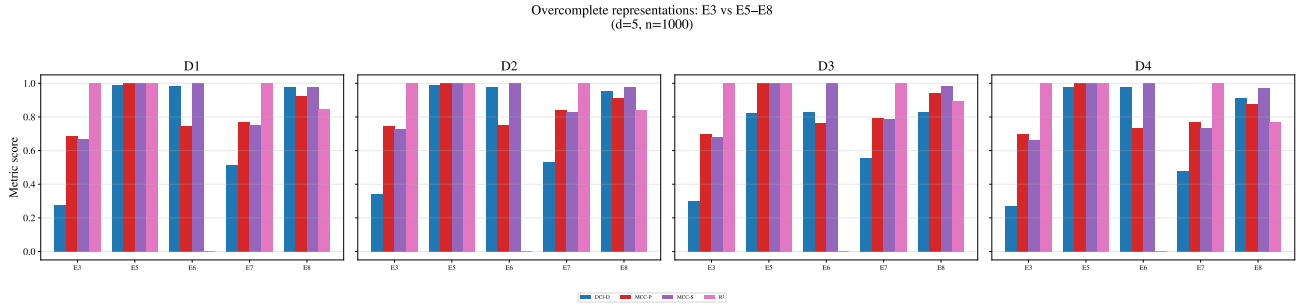


Figure 18: Overcomplete encoders across all DGP types ($d = 5, n = 1000$). Extension of Fig. 17 from $d = 20$ to $d = 5$. At this smaller d , DCI-D and MCC still separate disentangled overcomplete encoders (E5, E6, E8) from the entangled baseline (E3, E7), though the gap is narrower than at $d = 20$.

Factor predictability vs disentanglement
(coupling strength: ρ for D2, α for D3) (E1, $d=5$, $n=1000$)

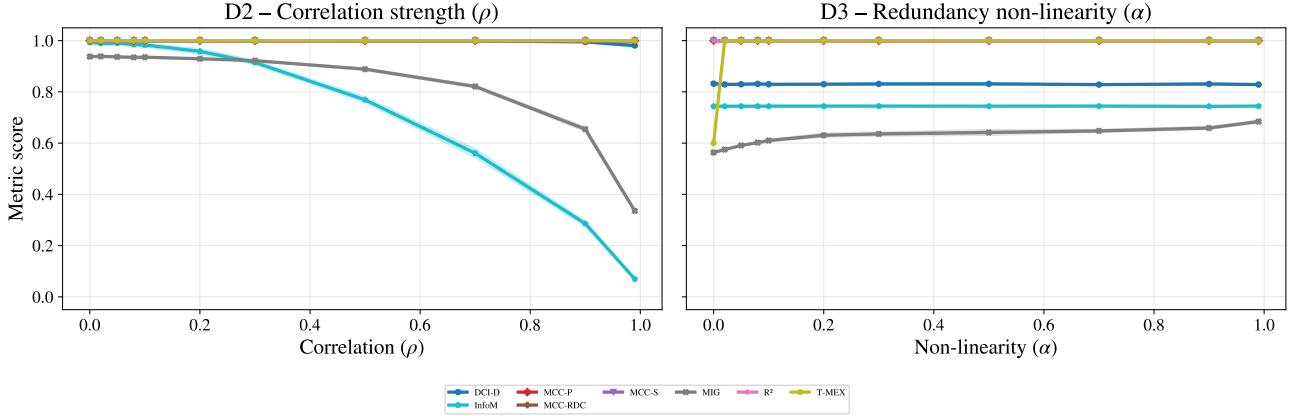


Figure 19: **Metrics conflate factor predictability with disentanglement under functional dependencies.** Full metric suite under **E1** comparing $\mathbf{D}_\rho$ (varying ρ) with $\mathbf{D}_f$ (elementwise nonlinear redundancy $z_2 = f(z_1)$). Under $\mathbf{D}_\rho$, only MIG and InfoM are affected by an increased correlation. Under $\mathbf{D}_f$, all metrics are almost invariant to nonlinear strength (excepted T-MEX, and MIG to a lesser extend), but DCI-D and InfoM do not reach 1.0 when there is such a redundancy. $d = 5$, $n = 1000$. See Fig. 20 for $d = 10$.

Factor predictability vs disentanglement
(coupling strength: ρ for D2, α for D3) (E1, $d=5$, $n=1000$)

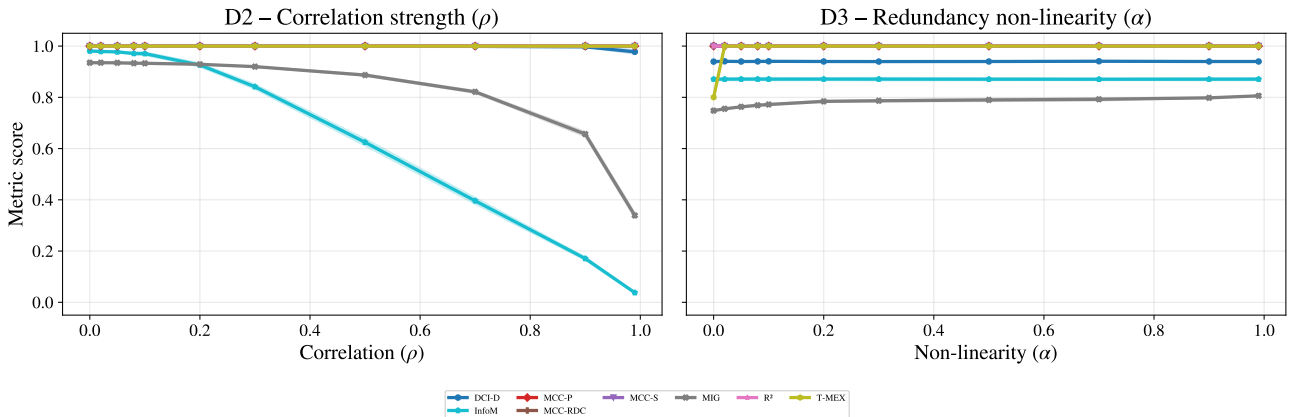


Figure 20: **Predictability vs. disentanglement at $d = 10$.** Same setup as Fig. 19 with $d = 10$ factors. The observations we made at $d = 5$ remain qualitatively valid at $d = 10$.

Redundancy (r) vs compression (m) – D3 + E4, $d=10$

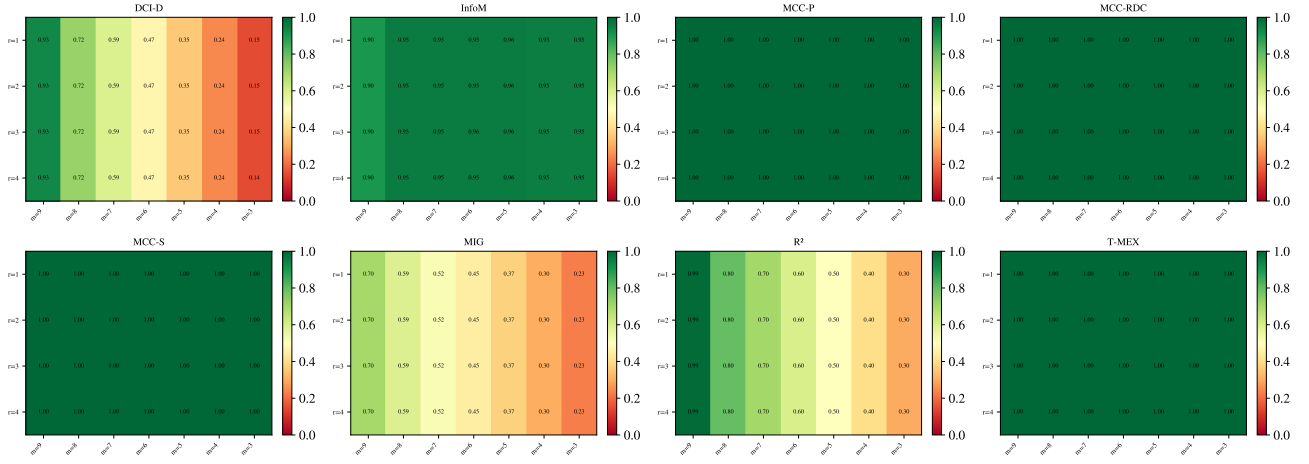


Figure 21: **Redundancy and compression under D_F .** The encoder compresses d factors (with a single redundant factor $z_2 = f(z_1)$, $d_{\text{eff}} = d-1$) into $m \leq d$ codes. R^2 and DCI-D plateau near 1.0 at $m = d_{\text{eff}}$, correctly recognising that the omitted factor carries no independent information. MCC-P/S report 1.0 at all compression levels, unable to distinguish lossless from lossy omission.

Effect of encoding type: E1 vs E5, E6, E8
($d=5$, $n=1000$)

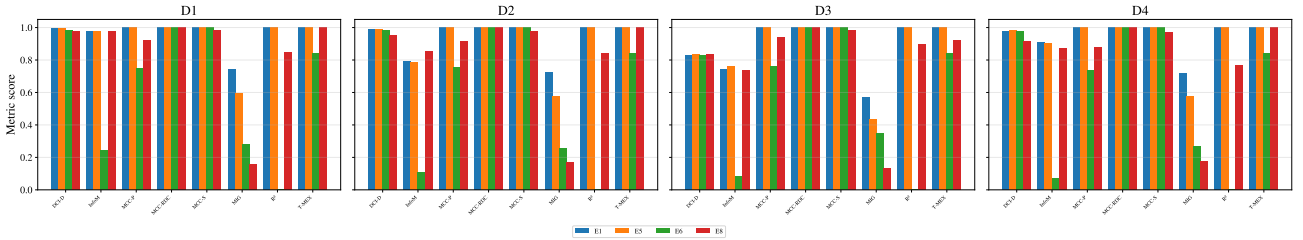


Figure 22: **Full metric suite: scores across encoding types and DGP types.** Each panel compares matched-dimension encoders (E1–E3) with overcomplete encoders (E5–E8) under $D_{\perp}$ – D_F . MI-based metrics (MIG, InfoMEC) and T-MEX are included alongside the main metrics. $d = 5$, $n = 1000$.

Overcomplete representations: E3 vs E5–E8
($d=5$, $n=1000$)

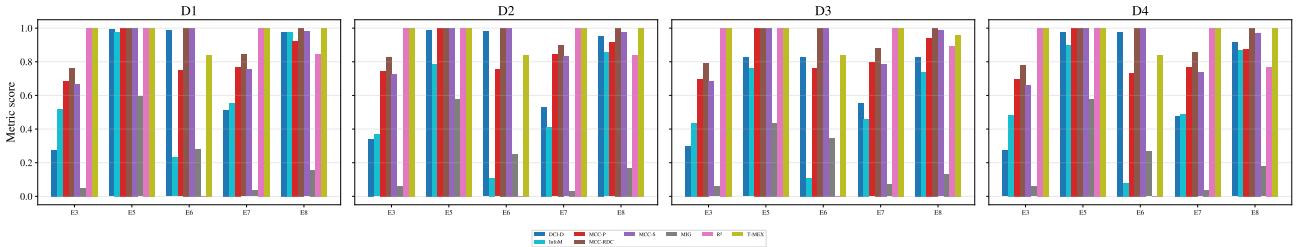


Figure 23: **Full metric suite for the overcomplete m/d sweep.** Extension of Fig. 4 to all metrics (MIG, InfoMEC, MCC-RDC, T-MEX). MI-based metrics decline for distributed codes (E8) similarly to MCC, while T-MEX is more robust to overcompleteness. $d = 5$, $n = 1000$.

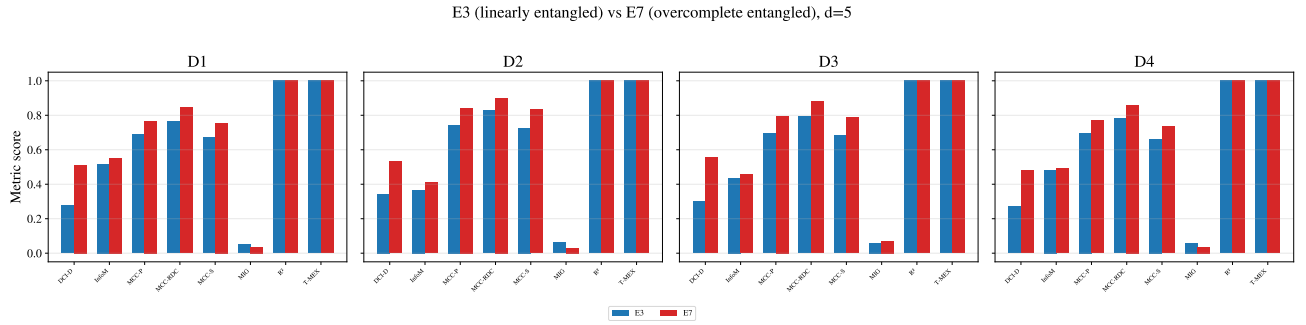


Figure 24: **Matched-dimension entangled (E3) versus overcomplete entangled (E7)**. Full metric suite comparing the two entangled geometries as m/d increases. DCI-D inflates for E7, while the rest of the methods return approximately equivalent scores. These observations are stable across DGPs. $d = 5$, $n = 1000$.

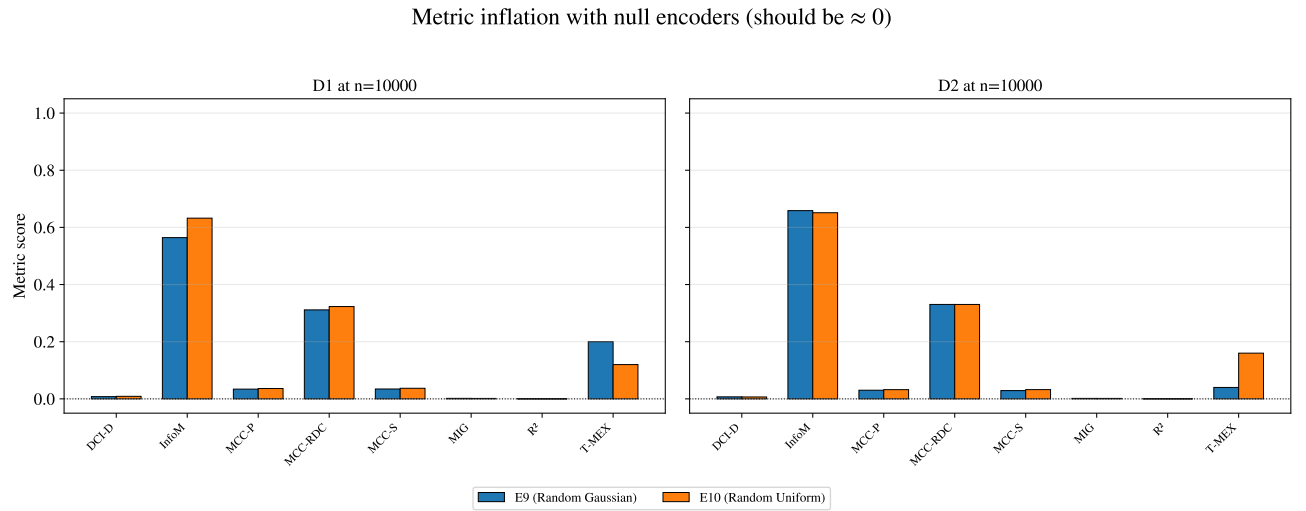


Figure 25: **Metric inflation under null encoders E9 and E10**. Full metric suite showing scores when the representation is independent of $\mathbf{z}$. MCC-RDC and InfoM exhibit persistent inflation that does not vanish at high n . R^2 , DCI-D, and MIG remain closest to the expected value of 0. $d = 5$, $n = 1000$.

Null-encoder convergence – do metrics reach 0?

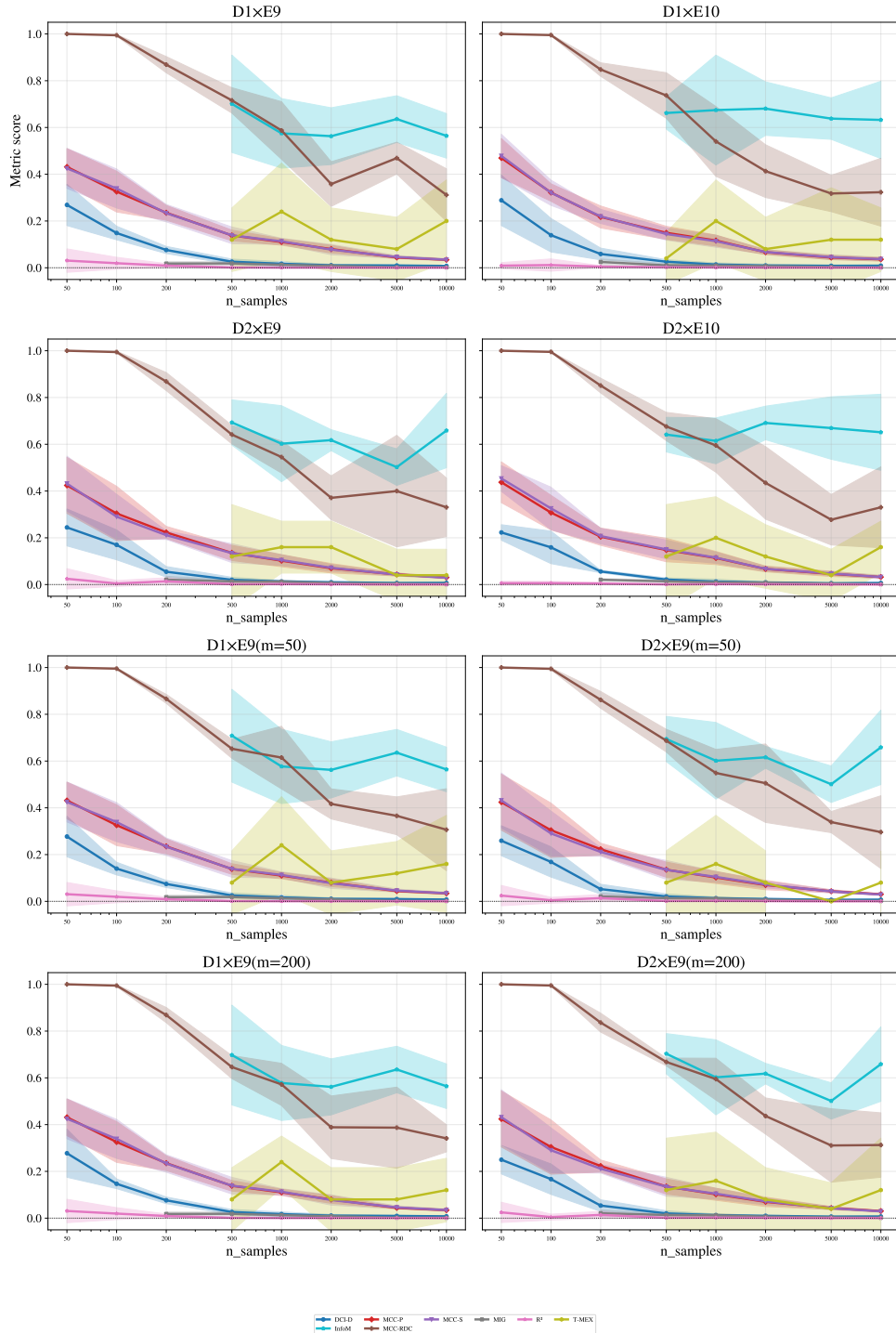


Figure 26: **Convergence of null-encoder scores with increasing n , for Gaussian or uniform distributions and varying m .** Both encoders **E9** and **E10** are independent from the ground truth factors, so scores should converge to 0 as n grows. We observe that R^2 converges fastest, followed by MIG and DCI-D. All variants of MCC converge more slowly, consistent with the $\sqrt{2 \log m/n}$ floor (§ G.3). Finally, InfoM does not seem to converge to 0 in the settings we tested. $d = 5$.

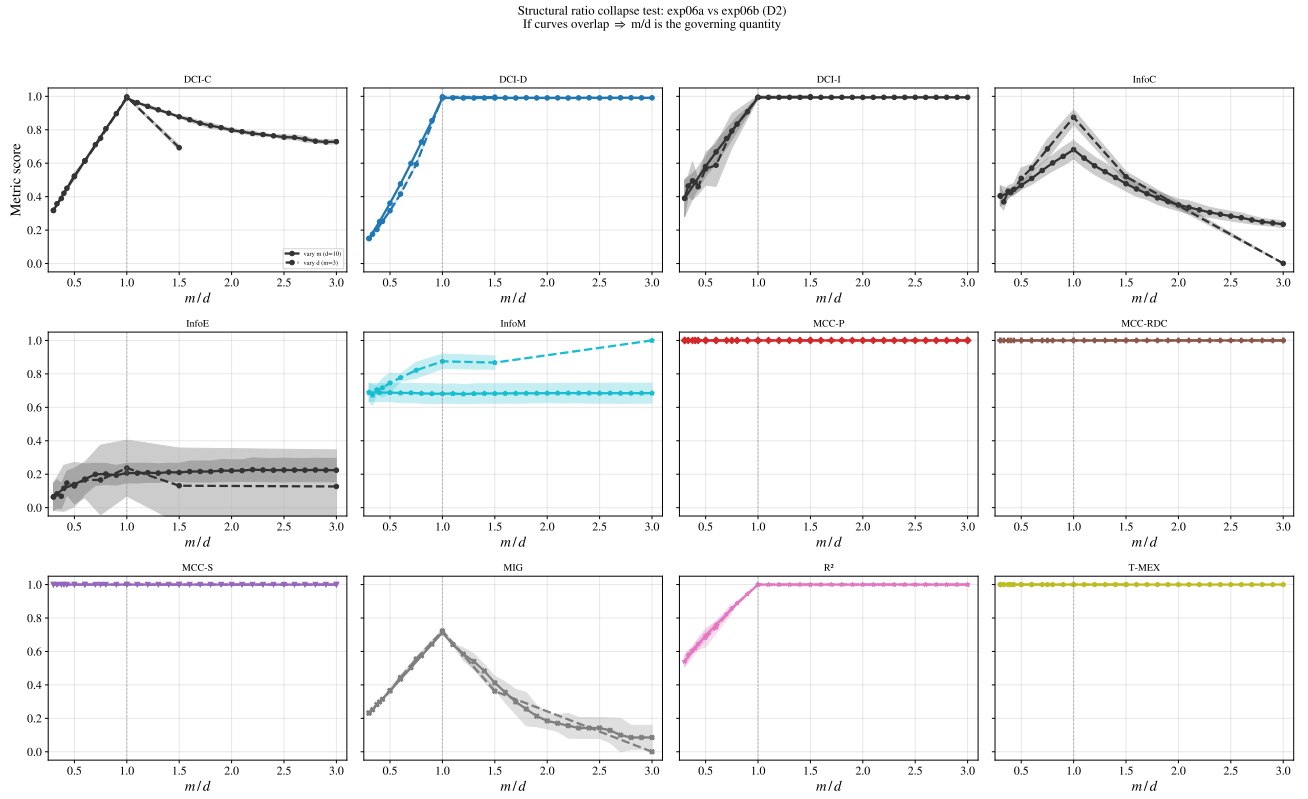


Figure 27: **Metric scores across encoder types under correlated factors (D_ρ)**. Overlay of all metrics for varying ρ under D_ρ . MCC-P/S increase with $|\rho|$ under entangled encoders (**E3**), while R^2 and DCI-D are less affected by the correlation structure. $d = 5$, $n = 1000$.

Ratio collapse test: metric score vs d/n ($D1 + E1$, $m = d$)
 Overlap $\Rightarrow d/n$ governs; separation $\Rightarrow d$ or n matters independently

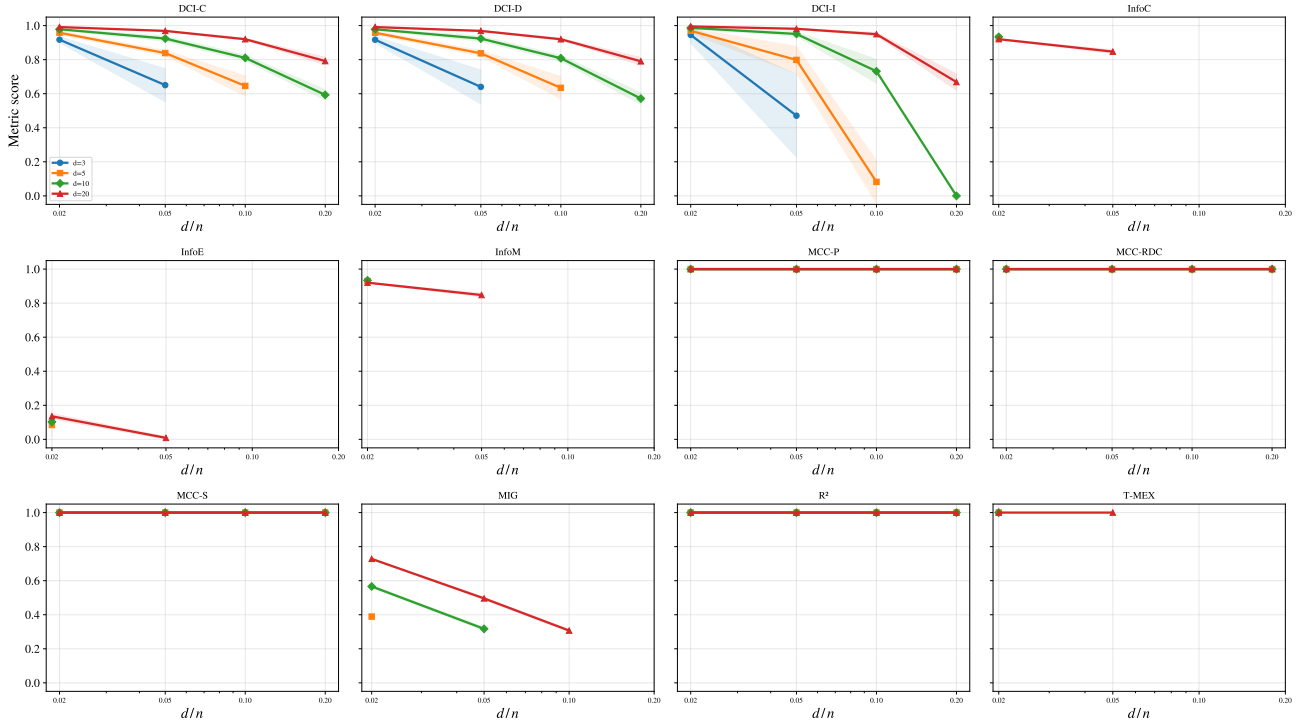


Figure 28: **Metric scores as a function of m/d under different (m, d) configurations.** Each panel shows one metric; overlapping curves from different (m, d) pairs with the same ratio confirm that m/d , not m or d individually, governs metric behaviour in the undercomplete regime. MCC-P/S report 1.0 regardless of m/d ; R^2 and DCI-D increase approximately linearly. $n = 1000$.

Metrics split into three ratio-dependence classes

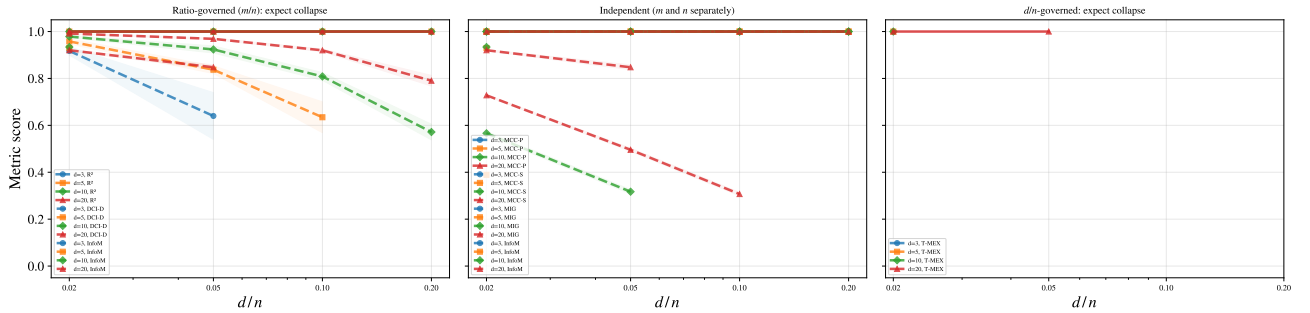


Figure 29: **Ratio-collapse analysis grouped by encoder type.** Same data as Fig. 28, reorganised by encoder geometry. Curves from sweeps over m (fixed d) and over d (fixed m) overlap at matched m/d , confirming the ratio as the governing quantity across encoder types.

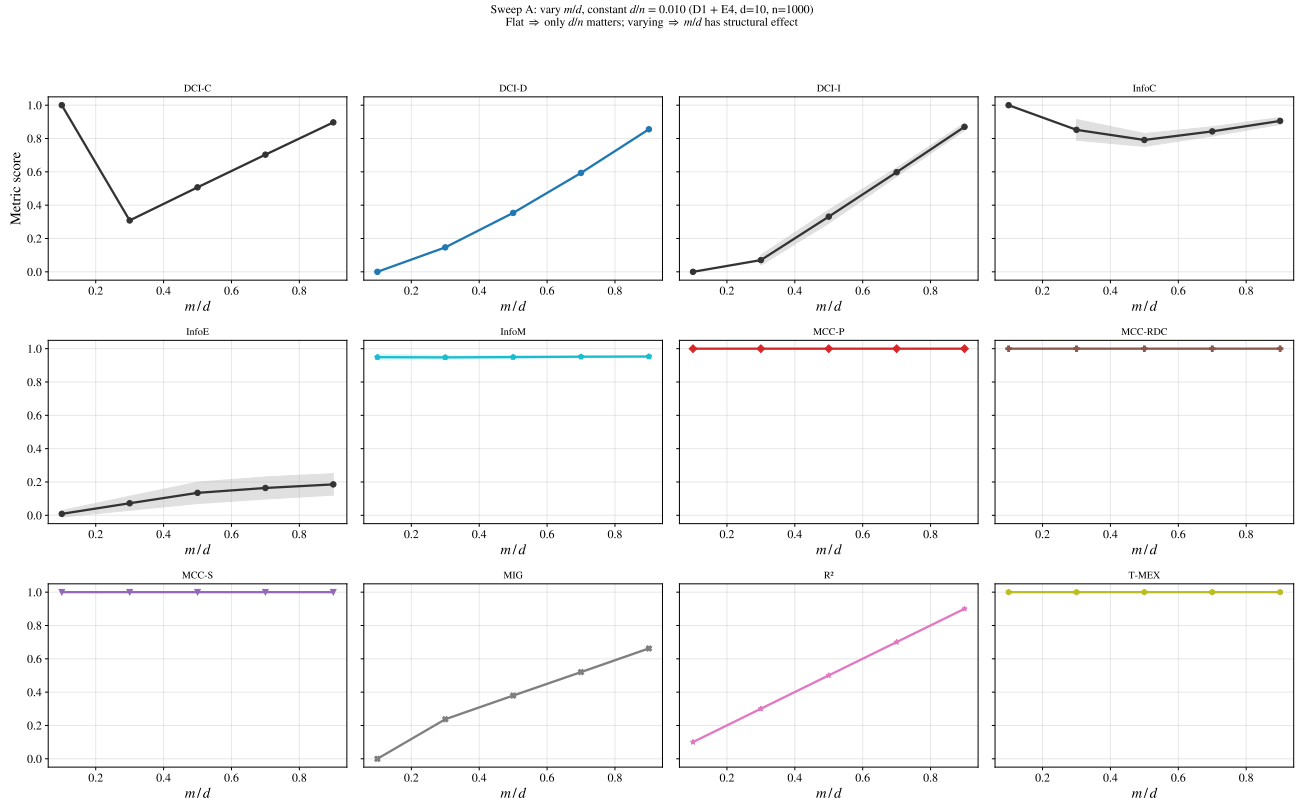


Figure 30: **Sensitivity to either d/n and m/d ratios.** Each metric is tested on the DGP $\mathbf{D}_{\perp}$ and encoder $\mathbf{E4}$ with varying m/d ratio, at constant $d/n = 0.01$. A flat curve means that only d/n matters; variations imply that m/d has a structural effect. All MCC variants, InfoM and T-MEX are insensitive to m/d , while the rest of the metrics show a sensitivity to it. $d = 5$, $n = 1000$. See Fig. 31 for the converse experiment.

Sweep B: vary d/n , constant $m/d = 0.50$ (D1 + E4, $d=10$, $m=5$)
Convergence study: metrics should improve as n grows

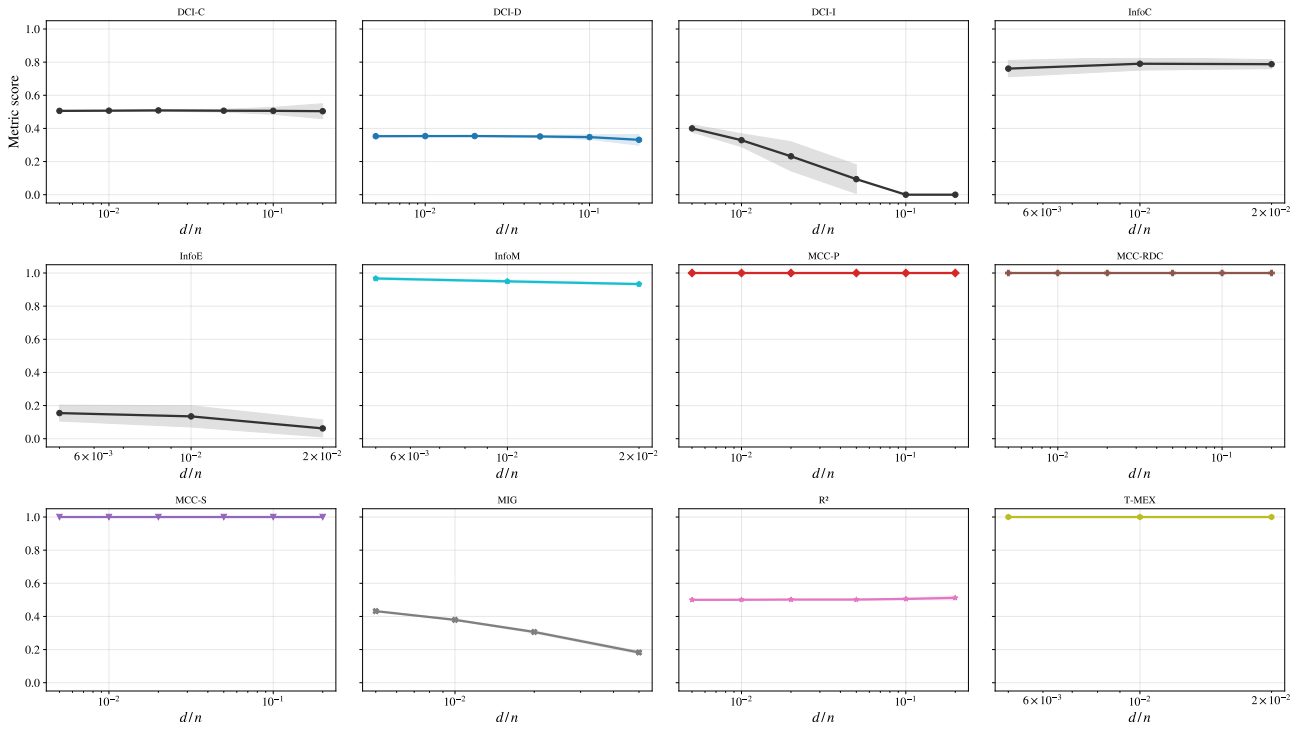


Figure 31: **Sensitivity to either d/n and m/d ratios.** Each metric is tested on the DGP $\mathbf{D}_{\perp}$ and encoder $\mathbf{E4}$ with varying d/n ratio, at constant $m/d = 0.5$. A flat curve means that only d/n matters; variations imply that m/d has a structural effect. All metrics are almost flat, except MI-based ones (in particular, MIG, InfoE, InfoC). See Fig. 30 for the converse experiment.

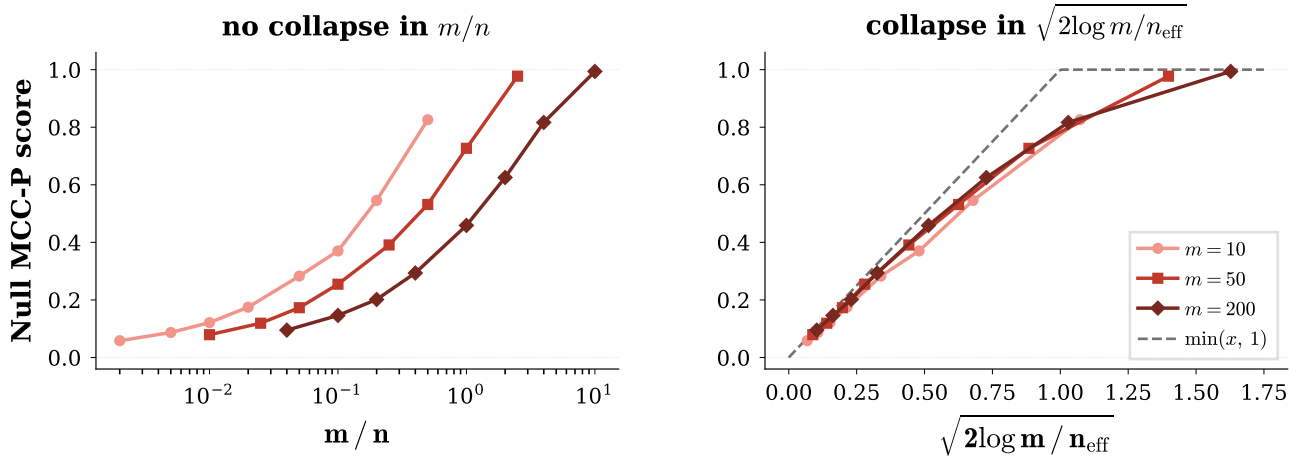


Figure 32: **Null MCC-P inflation collapses onto the extreme-value floor $\sqrt{2 \log m / n_{\text{eff}}}$.** Null-encoder MCC-P for $m \in \{10, 50, 200\}$ swept over test-split size. *Left:* plotted against m/n , the curves for different m do *not* align—the false-positive floor is not a function of m/n alone. *Right:* plotted against $\sqrt{2 \log m / n_{\text{eff}}}$, all three curves collapse onto a single master curve tracking $\min(x, 1)$ (dashed), confirming $\mathbb{E}[\text{MCC-P}] \gtrsim \sqrt{2 \log m / n}$ derived in § G.3. n_{eff} is the held-out split (20% of n) on which the metric is scored.

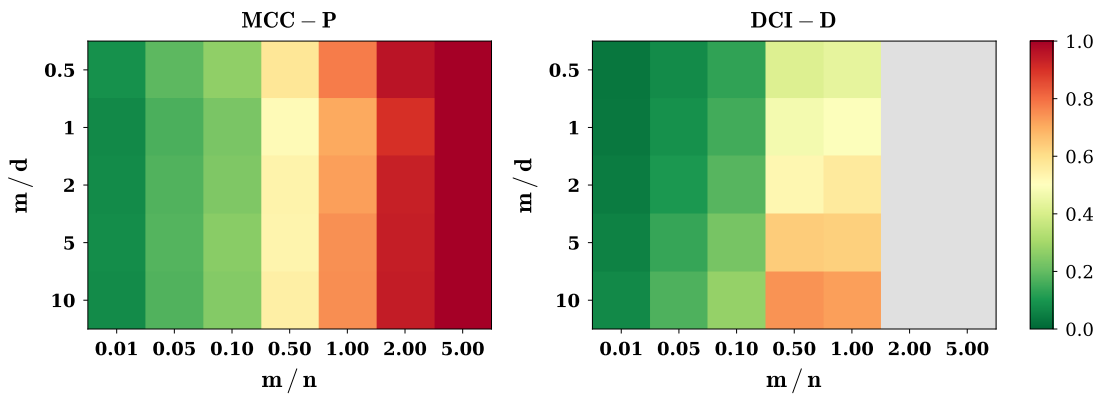


Figure 33: False positive inflation under E10.

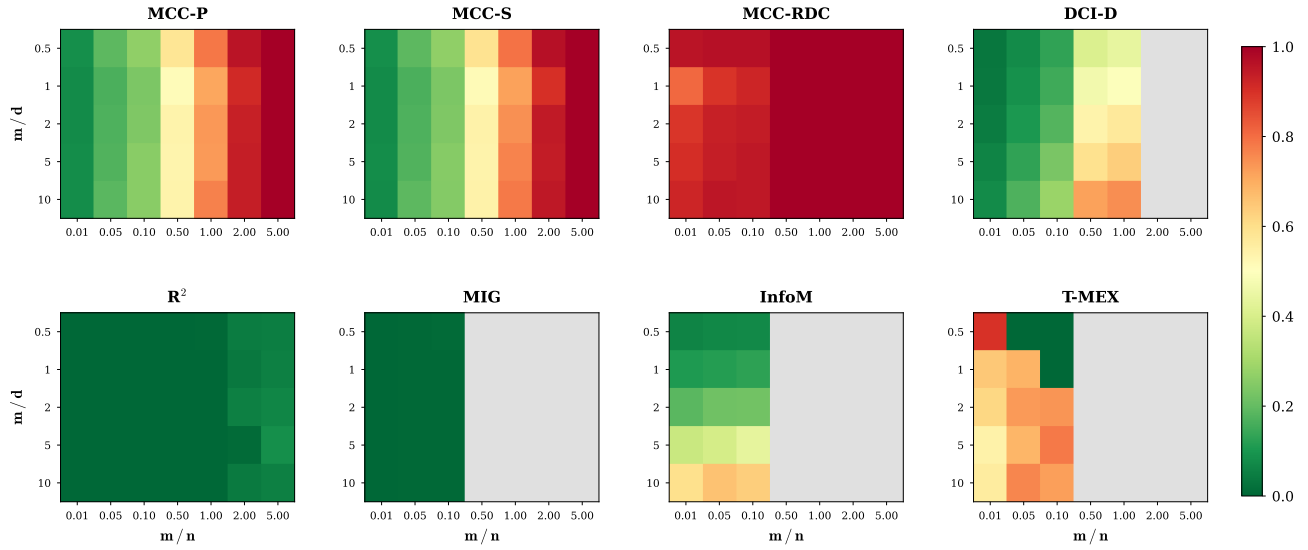


Figure 34: **All-metric false-positive phase diagram under the uniform null E10.** Full-suite extension of Fig. 5: each cell is a metric’s score for a random encoder independent of $\mathbf{z}$, so any value > 0 is a false positive. MCC-P/S/RDC inflate with m/n and are flat in m/d (worst: MCC-RDC, inflated almost everywhere), whereas R^2 and MIG stay ≈ 0 . Grey cells are $(m/d, m/n)$ points where a metric is undefined at small n (MIG, InfoM, T-MEX).

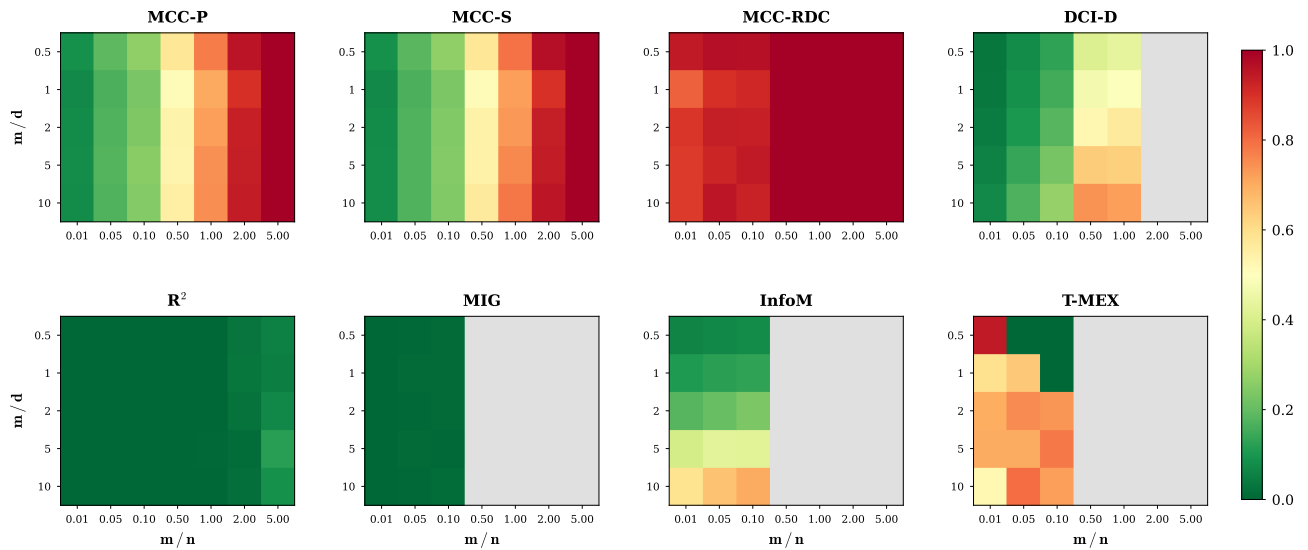


Figure 35: **All-metric false-positive phase diagram under the Gaussian null E9.** We observe a similar metric inflation issue as with E10.