
Contrastive MIM: A Contrastive Mutual Information Framework for Unified Generative and Discriminative Representation Learning

Anonymous Author(s)

Affiliation

Address

email

Abstract

We present *Contrastive Mutual Information Machine* (cMIM), a probabilistic framework that adds a contrastive objective to the Mutual Information Machine (MIM) Livne et al. [2019], yielding representations that are effective for both discriminative and generative tasks. Unlike conventional contrastive learning van den Oord et al. [2018], Chen et al. [2020], He et al. [2020], cMIM does not require positive data augmentations and exhibits reduced sensitivity to batch size. We further introduce *informative embeddings*, a generic training-free method to extract enriched features from decoder hidden states of encoder-decoder models.

We evaluate cMIM on life-science tasks, including molecular property prediction on ZINC-15 Sterling and Irwin [2015], Weininger [1988] (ESOL, FreeSolv, Lipophilicity) and biomedical image classification (MedMNIST Yang et al. [2021]). cMIM consistently improves downstream accuracy over MIM and InfoNCE baselines while maintaining comparable reconstruction quality. These results indicate cMIM is a promising foundation-style representation learner for biomolecular and biomedical applications and is readily extendable to multi-modal settings (e.g., molecules + omics + imaging).

1 Introduction

Learning representations that support both unknown downstream prediction tasks and generative use cases is central to life-science machine learning, where data span molecules, sequences, imaging, and multi-omics. Contrastive losses such as InfoNCE van den Oord et al. [2018], Chen et al. [2020] are effective but often hinge on meaningful augmentations and large batches. Probabilistic auto-encoders maximize information about inputs, yet their latent geometry can be suboptimal for discriminative tasks Livne et al. [2019], Reidenbach et al. [2023].

We propose **cMIM**, a simple contrastive extension of MIM that (i) introduces a discriminator over pair-similarity without requiring positive augmentations, (ii) aligns local clustering with global angular separation, and (iii) remains robust to batch-size via Monte Carlo expectations Hoeffding [1963]. We also propose **informative embeddings**, obtained from decoder hidden states, that substantially improve downstream performance *without extra training*.

We evaluate cMIM on molecular property prediction on ZINC-15 Sterling and Irwin [2015], Weininger [1988] (ESOL, FreeSolv, Lipophilicity) and biomedical imaging with *MedMNIST* Yang et al. [2021]. Experiments show consistent gains over MIM and InfoNCE with reduced batch-size sensitivity.

Algorithm 1 Learning parameters θ of cMIM

Require: Samples from dataset $\mathcal{P}(\mathbf{x})$

```

1: while not converged do
2:    $\mathcal{D} \leftarrow \{\mathbf{x}_j, \mathbf{z}_j \sim q_\theta(\mathbf{z} | \mathbf{x}) \mathcal{P}(\mathbf{x})\}_{j=1}^B$  {Sample a batch}
3:    $\hat{\mathcal{L}}_{\text{A-MIM}} = -\frac{1}{B} \sum_{i=1}^B (\log p_\theta(\mathbf{x}_i | \mathbf{z}_i) + \log p_{k=1}(\mathbf{x}_i, \mathbf{z}_i) + \frac{1}{2}(\log q_\theta(\mathbf{z}_i | \mathbf{x}_i) + \log p(\mathbf{z}_i)))$ 
4:    $\theta \leftarrow \theta - \eta \nabla_\theta \hat{\mathcal{L}}_{\text{A-MIM}}$  {Reparameterized gradients}
5: end while
  
```

Figure 1: Training algorithm for cMIM.

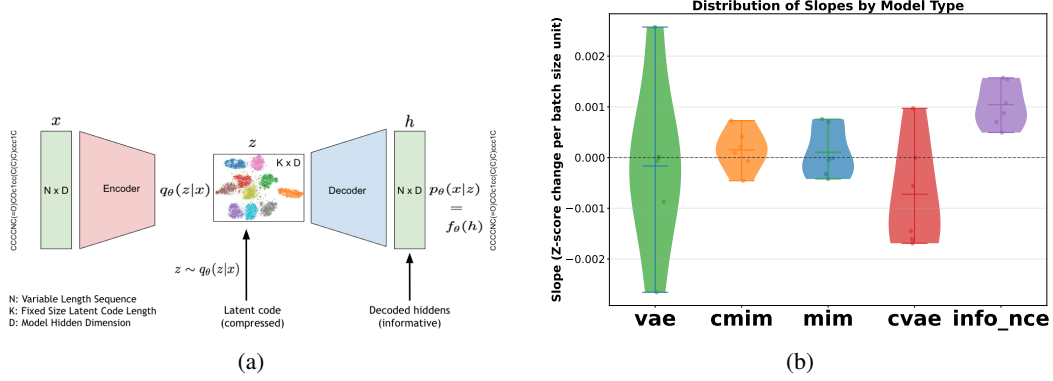


Figure 2: Left: **(a)** Informative embeddings $\mathbf{h}$ are taken from decoder hidden states before mapping to $p_\theta(\mathbf{x} | \mathbf{z})$. For autoregressive decoders we use teacher forcing. Right: **(b)** Distribution of slopes from linear fits of accuracy vs. batch-size for different models. Each point corresponds to the average z-score of a model trained on MNIST-like datasets. Both MIM and cMIM are not sensitive to batch-size.

2 Formulation

Preliminaries. Let $\mathbf{x}$ denote the observation and $\mathbf{z}$ the latent code. MIM is a probabilistic auto-encoder that maximizes mutual information between $\mathbf{x}$ and $\mathbf{z}$ while encouraging clustered latents via entropy minimization.

Contrastive variable and objective. We introduce a binary variable k that indicates whether $(\mathbf{x}, \mathbf{z})$ is a matched pair ($k=1$) or a mismatched pair ($k=0$). A matched pair is defined as $\mathbf{z}_i \sim q_\theta(\mathbf{z} | \mathbf{x}_i)$. We define encoder/decoder factorizations with a discriminator over k :

$$q_\theta(\mathbf{x}, \mathbf{z}, k) = q_\theta(k | \mathbf{x}, \mathbf{z}) q_\theta(\mathbf{z} | \mathbf{x}) q_\theta(\mathbf{x}), \quad p_\theta(\mathbf{x}, \mathbf{z}, k) = p_\theta(k | \mathbf{x}, \mathbf{z}) p_\theta(\mathbf{x} | \mathbf{z}) p_\theta(\mathbf{z}). \quad (1)$$

Let $\mathbf{z}_i \sim q_\theta(\mathbf{z} | \mathbf{x}_i)$. Using a temperature-scaled cosine similarity $s(\mathbf{z}_i, \mathbf{z}_j)/\tau$ with $g_{ij} \equiv g(\mathbf{z}_i, \mathbf{z}_j) = \exp(s(\mathbf{z}_i, \mathbf{z}_j)/\tau)$, we set

$$p_{k=1}(\mathbf{x}_i, \mathbf{z}_i) = \frac{g_{ii}}{g(\mathbf{z}_i, \mathbf{z}_i) + \mathbb{E}_{\mathbf{x}' \sim P(\mathbf{x}), \mathbf{z}' \sim q_\theta(\mathbf{z} | \mathbf{x}')} [g(\mathbf{z}_i, \mathbf{z}')] } \approx \frac{g_{ii}}{g_{ii} + \frac{1}{B-1} \sum_{j \neq i} g_{ij}}. \quad (2)$$

Training samples always satisfy $k=1$ since $\mathbf{z}_i$ is drawn conditionally on $\mathbf{x}_i$. Replacing the denominator expectation by an in-batch Monte Carlo estimate yields a simple, augmentation-free contrastive term with concentration improving in B .

Relation to InfoNCE. Defining $s_{ij} = s(\mathbf{z}_i, \mathbf{z}_j)/\tau$ makes Eq. (2) equivalent to an InfoNCE softmax where the positive logit is offset by $\log(B-1)$. This calibrates $p_{k=1}$ to $1/2$ when logits are equal (independent of B), reducing batch-size sensitivity. See Appendix for more details.

Objective. We adopt the A-MIM upper bound [Livne et al., 2020] with the added contrastive term:

$$\hat{\mathcal{L}}_{\text{A-MIM}}(\theta; \mathcal{D}) = -\frac{1}{N} \sum_{i=1}^N \left[\log p_\theta(\mathbf{x}_i | \mathbf{z}_i) + \log p_{k=1}(\mathbf{x}_i, \mathbf{z}_i) + \frac{1}{2} \left(\log q_\theta(\mathbf{z}_i | \mathbf{x}_i) + \log P(\mathbf{z}_i) \right) \right], \quad (3)$$

where $P(\mathbf{z})$ is a Normal anchor prior. The first term preserves generative fidelity; the second encourages global angular separation.

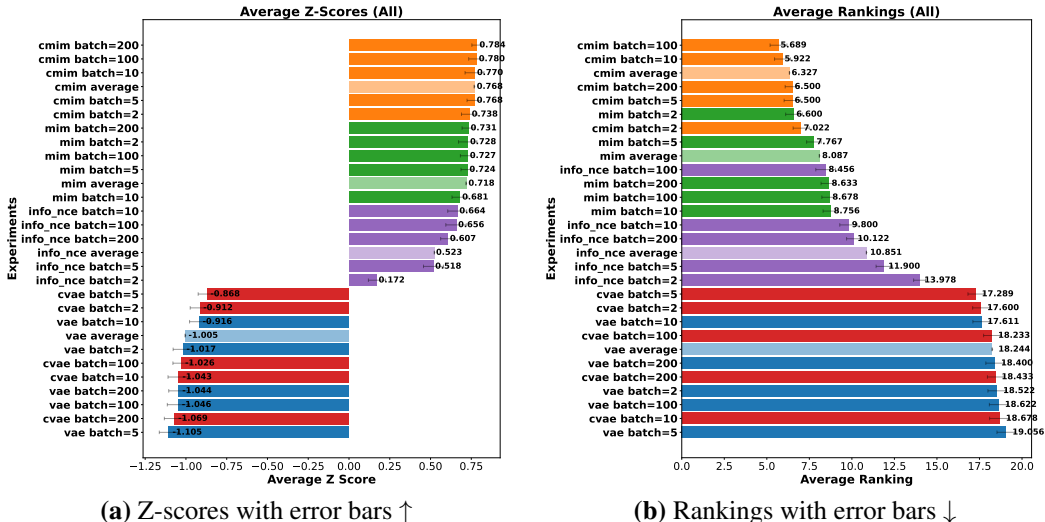


Figure 3: Classification accuracy across datasets and classifiers. Colors indicate model families: cMIM (orange), MIM (green), InfoNCE (purple), VAE (blue), cVAE (red). Light shades denote model averages. cMIM consistently outperforms all baselines across batch-sizes and metrics.

Model (Latent $K \times D$)	ESOL		FreeSolv		Lipophilicity		Recon.
	SVM	MLP	SVM	MLP	SVM	MLP	
MIM (1×512) [reproduction]	0.65	0.34	2.23	1.82	0.663	0.61	100%
cMIM (1×512)	0.47	0.19	2.32	1.67	0.546	0.38	100%
MIM (1×512) info emb	0.21	0.29	1.55	1.4	0.234	0.28	100%
cMIM (1×512) info emb	0.21	0.24	1.74	1.35	0.24	0.23	100%
CDDD (512)	0.33		0.94		0.4		
$\dagger$ Seq2seq ($N \times 512$)	0.37	0.43	1.24	1.4	0.46	0.61	100%
$\dagger$ Perceiver (4×512)	0.4	0.36	1.22	1.05	0.48	0.47	100%
$\dagger$ VAE (4×512)	0.55	0.49	1.65	3.3	0.63	0.55	46%
MIM (1×512)	0.58	0.54	1.95	1.9	0.66	0.62	100%
Morgan fingerprints (512)	1.52	1.26	5.09	3.94	0.63	0.61	

Table 1: Comparison of models on ESOL, FreeSolv, and Lipophilicity using SVM and MLP regressors, with reconstruction accuracy. Top: our results. Bottom: results from Reidenbach et al. [2023]. For $\dagger$ models, sequence representations were averaged to 512 dimensions. Bold: best non-MIM results. Highlighted: best among MIM-based models. Note that CDDD training included these classification tasks.

Informative embeddings. Instead of using z directly, we take the decoder hidden states h before mapping to $p_\theta(x|z)$ parameters, and pool them (mean over sequence/space) to obtain h_i . For autoregressive decoders we use teacher forcing, i.e., $h_i = \text{Decoder}(x_i, z_i)$. These *informative embeddings* enrich z with a context of the decoder distribution at no extra training cost. See Figure 2a.

3 Experiments

Setup. We report biomolecular and biomedical classification results.

Molecular property prediction. Following Reidenbach et al. [2023], we train on a large tranche of ZINC-15 Sterling and Irwin [2015] with SMILES tokenization Weininger [1988] and evaluate ESOL, FreeSolv, and Lipophilicity on held-out splits. We compare MIM Livne et al. [2019], cMIM, VAE Kingma and Ba [2015], and an InfoNCE encoder van den Oord et al. [2018]. Embeddings are

61 either (i) the mean encoder code z , or (ii) *informative embeddings* h from the decoder. Downstream
62 regressors are SVM/MLP (Scikit-learn defaults Pedregosa et al. [2011]). See Table 1.

63 **Biomedical imaging.** We evaluate on MedMNIST Yang et al. [2021] as a light-weight benchmark
64 for learning transferable biomedical image representations. All models share the same Perceiver-style
65 encoder/decoder Jaegle et al. [2021]; InfoNCE uses the encoder only. We report accuracy across
66 datasets and summarize with average ranks/z-scores. See Figure 3.

67 **Key Results.** (i) *Biomolecules.* cMIM consistently improves ESOL/FreeSolv/Lipophilicity errors
68 relative to MIM and VAE when using informative embeddings, and is competitive with chemical
69 baselines (e.g., CDDD Winter et al. [2019]) while retaining perfect reconstruction.

70 (ii) *Biomedical imaging.* Across MedMNIST tasks, cMIM dominates average rankings over
71 MIM/InfoNCE/VAE across batch-sizes.

72 (iii) *Robustness.* Slopes from accuracy vs. batch-size fits cluster near zero for cMIM and MIM, while
73 InfoNCE increases with batch-size, confirming reduced sensitivity. cMIM, like MIM, remains stable
74 even with very small batch-sizes, unlike InfoNCE which degrades as batch-size decreases.

75 (iv) *Generative fidelity.* cMIM matches MIM reconstruction on molecules and shows slight improve-
76 ments on biomedical images (Appendix Figure 6), suggesting a benign regularization effect.

77 4 Related Work

78 **Contrastive learning.** CPC/InfoNCE van den Oord et al. [2018], SimCLR Chen et al. [2020], and
79 MoCo He et al. [2020] demonstrated strong representation learning but rely on augmentations and
80 many negatives. Augmentation-free variants such as BYOL Grill et al. [2020] or SimSiam Chen and
81 He [2021] introduce architectural asymmetries/predictors.

82 **Mutual-information auto-encoding.** MIM Livne et al. [2019] maximizes mutual information
83 while clustering latents; related MI-regularized VAEs/auto-encoders optimize information-theoretic
84 surrogates but often require intricate weighting. Our contribution integrates a calibrated contrastive
85 term into a probabilistic MIM, improving global discriminative geometry with minimal complexity.

86 **Life-science foundation modeling.** Sequence-to-sequence and continuous descriptor models (e.g.,
87 CDDD Winter et al. [2019]) support molecular property prediction; MedMNIST Yang et al. [2021]
88 popularizes light-weight biomedical imaging benchmarks. Our results show cMIM’s unified gener-
89 ative+discriminative modeling is competitive across both, and the method is directly extensible to
90 multi-modal use cases.

91 5 Limitations

92 While cMIM improves discriminative performance and batch-size robustness, limitations remain.
93 Generative evaluation is restricted to reconstruction, leaving open questions on sample quality
94 and controllability. Scalability to larger models and modalities (e.g., video, long-context text) is
95 untested. Performance may still depend on the similarity function, temperature τ , and the number of
96 negatives, which adds computational cost. Future work should address scaling, broader modalities,
97 and generative analysis.

98 6 Conclusions

99 We introduced cMIM, a simple augmentation-free contrastive extension to probabilistic represen-
100 tation learning that improves discriminative performance while maintaining generative fidelity. On
101 molecular properties Sterling and Irwin [2015], Weininger [1988], Winter et al. [2019] and biomed-
102 ical imaging Yang et al. [2021], cMIM (especially with informative embeddings) outperforms
103 MIM/InfoNCE/VAEs and is robust to batch-size.

104 Specifically, cMIM offers a practical backbone for life-science foundation models: its factoriza-
105 tion and decoder-based embeddings extend naturally to multi-modal molecular+omics or molecu-
106 lar+imaging setups. Future work: (i) multi-modal joint training, (ii) scaling to larger architectures
107 Jaegle et al. [2021] and datasets, and (iii) systematic studies of geometry and calibration Wang and
108 Isola [2020].

References

- Alexander Alemi, Ben Poole, Ian Fischer, Joshua Dillon, Rif A. Saurous, and Kevin Murphy. Fixing a broken elbow. In *International Conference on Machine Learning (ICML)*, pages 159–168, 2018.
- Steven Bird, Ewan Klein, and Edward Loper. *Natural Language Processing with Python: Analyzing Text with the Natural Language Toolkit*. O’Reilly Media, Inc., 2009.
- Ting Chen, Simon Kornblith, Mohammad Norouzi, and Geoffrey E. Hinton. A simple framework for contrastive learning of visual representations. *arXiv preprint arXiv:2002.05709*, 2020.
- Xinlei Chen and Kaiming He. Exploring simple siamese representation learning. *arXiv preprint arXiv:2011.10566*, 2021.
- Gregory Cohen, Saeed Afshar, Jonathan Tapson, and André van Schaik. EMNIST: an extension of MNIST to handwritten letters. *CoRR*, abs/1702.05373, 2017. URL <http://arxiv.org/abs/1702.05373>.
- Li Deng. The mnist database of handwritten digit images for machine learning research [best of the web]. *IEEE Signal Processing Magazine*, 29(6):141–142, 2012. doi: 10.1109/MSP.2012.2211477.
- Jean-Bastien Grill, Florian Strub, Florent Altché, Corentin Tallec, Pierre H. Richemond, Elena Buchatskaya, Carl Doersch, Bernardo Ávila Pires, Zhaohan Daniel Guo, Mohammad Gheshlaghi Azar, Bilal Piot, Koray Kavukcuoglu, Rémi Munos, and Michal Valko. Bootstrap your own latent: A new approach to self-supervised learning. *arXiv preprint arXiv:2006.07733*, 2020.
- Kaiming He, Haoqi Fan, Yuxin Wu, Saining Xie, and Ross Girshick. Momentum contrast for unsupervised visual representation learning. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 9729–9738, 2020.
- Irina Higgins, Loic Matthey, Arka Pal, Christopher Burgess, Xavier Glorot, Matthew Botvinick, Shakir Mohamed, and Alexander Lerchner. β -VAE: Learning Basic Visual Concepts with a Constrained Variational Framework. In *International Conference on Learning Representations (ICLR)*, 2017.
- Wassily Hoeffding. Probability inequalities for sums of bounded random variables. *Journal of the American Statistical Association*, 58(301):13–30, March 1963. URL <http://www.jstor.org/stable/2282952>.
- Shengding Hu, Yuge Tu, Xu Han, Chaoqun He, Ganqu Cui, Xiang Long, Zhi Zheng, Yewei Fang, Yuxiang Huang, Weilin Zhao, et al. Minicpm: Unveiling the potential of small language models with scalable training strategies. *arXiv preprint arXiv:2404.06395*, 2024.
- Ross Irwin, Spyridon Dimitriadis, Jiazhen He, and Esben Jannik Bjerrum. Chemformer: a pre-trained transformer for computational chemistry. *Machine Learning: Science and Technology*, 3(1):015022, 2022. doi: 10.1088/2632-2153/ac3ffb. URL <https://doi.org/10.1088/2632-2153/ac3ffb>.
- Andrew Jaegle, Felix Gimeno, Andy Brock, Oriol Vinyals, Andrew Zisserman, and Joao Carreira. Perceiver: General perception with iterative attention. In *International Conference on Machine Learning (ICML)*, volume 139 of *Proceedings of Machine Learning Research*, pages 4651–4664. PMLR, 2021.
- Sunghwan Kim, Jie Chen, Tiejun Cheng, Asta Gindulyte, Jia He, Siqian He, Qingliang Li, Benjamin A. Shoemaker, Paul A. Thiessen, Bo Yu, Leonid Zaslavsky, Jian Zhang, and Evan E. Bolton. Pubchem 2019 update: improved access to chemical data. *Nucleic Acids Research*, 47(D1):D1102–D1109, 2018. doi: 10.1093/nar/gky1033. URL <https://doi.org/10.1093/nar/gky1033>.
- Diederik P. Kingma and Jimmy Ba. Adam: A method for stochastic optimization. In *International Conference on Learning Representations (ICLR)*, 2015.
- Oleksii Kuchaiev, Jason Li, Huyen Nguyen, Oleksii Hrinchuk, Ryan Leary, Boris Ginsburg, Samuel Kriman, Stanislav Beliaev, Vitaly Lavrukhin, Jack Cook, et al. Nemo: a toolkit for building ai applications using neural modules. *arXiv preprint arXiv:1909.09577*, 2019.

- 156 Micha Livne, Kevin Swersky, and David J. Fleet. MIM: Mutual Information Machine. *arXiv preprint*
157 *arXiv:1910.03175*, 2019.
- 158 Micha Livne, Kevin Swersky, and David J Fleet. Sentencemim: A latent variable language model.
159 *arXiv preprint arXiv:2003.02645*, 2020.
- 160 Fabian Pedregosa, Gaël Varoquaux, Alexandre Gramfort, Vincent Michel, Bertrand Thirion, Olivier
161 Grisel, Mathieu Blondel, Peter Prettenhofer, Ron Weiss, Vincent Dubourg, Jake Vanderplas,
162 Alexandre Passos, David Cournapeau, Matthieu Brucher, Matthieu Perrot, and Édouard Duchesnay.
163 Scikit-learn: Machine learning in python. *Journal of Machine Learning Research*, 12(85):2825–
164 2830, 2011. URL <http://jmlr.org/papers/v12/pedregosa11a.html>.
- 165 Danny Reidenbach, Micha Livne, Rajesh K. Ilango, Michelle Gill, and Johnny Israeli. Improving
166 small molecule generation using mutual information machine. *arXiv preprint arXiv:2208.09016*,
167 2023.
- 168 Teague Sterling and John J. Irwin. Zinc 15 – ligand discovery for everyone. *Journal of Chemical*
169 *Information and Modeling*, 55(11):2324–2337, 2015. doi: 10.1021/acs.jcim.5b00559. URL
170 <https://doi.org/10.1021/acs.jcim.5b00559>.
- 171 Aäron van den Oord, Yazhe Li, and Oriol Vinyals. Representation learning with contrastive predictive
172 coding. *arXiv preprint arXiv:1807.03748*, 2018.
- 173 Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N. Gomez,
174 Łukasz Kaiser, and Illia Polosukhin. Attention is all you need. In *Advances in Neural Information*
175 *Processing Systems (NeurIPS)*, volume 30, 2017.
- 176 Tongzhou Wang and Phillip Isola. Understanding contrastive representation learning through align-
177 ment and uniformity on the hypersphere. In *Proceedings of the 37th International Conference on*
178 *Machine Learning (ICML)*, pages 9929–9939. PMLR, 2020.
- 179 David Weininger. Smiles, a chemical language and information system. 1. introduction to methodol-
180 ogy and encoding rules. *Journal of Chemical Information and Computer Sciences*, 28(1):31–36,
181 1988. doi: 10.1021/ci00057a005. URL <https://doi.org/10.1021/ci00057a005>.
- 182 Robin Winter, Floriane Montanari, Frank Noé, and Djork-Arné Clevert. Learning continuous and
183 data-driven molecular descriptors by translating equivalent chemical representations. *Chemical*
184 *Science*, 10:1692–1701, 2019. doi: 10.1039/C8SC04175J. URL [http://dx.doi.org/10.1039/](http://dx.doi.org/10.1039/C8SC04175J)
185 [C8SC04175J](http://dx.doi.org/10.1039/C8SC04175J).
- 186 Jiancheng Yang, Rui Shi, Donglai Wei, Zequan Liu, Lin Zhao, Bilian Ke, Hanspeter Pfister, and
187 Bingbing Ni. Medmnist v2: A large-scale lightweight benchmark for 2d and 3d biomedical image
188 classification. *CoRR*, abs/2110.14795, 2021. URL <https://arxiv.org/abs/2110.14795>.

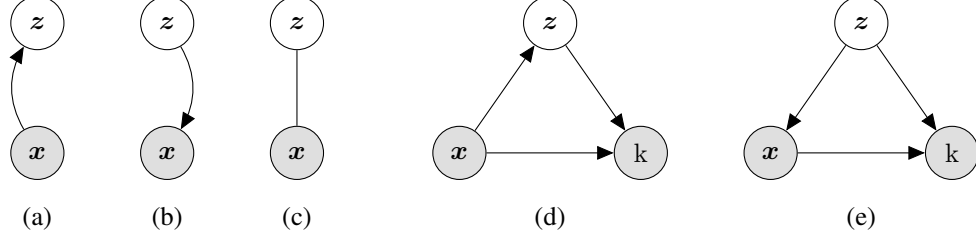


Figure 4: (Left) A MIM model learns two factorizations of a joint distribution (encoding/decoding) and an undirected joint. (Right) cMIM extends MIM with a binary variable k to encourage global discriminative structure while preserving local clustering.

A Extended Formulation

Background: Contrastive Learning

Intuition. Contrastive learning pulls positives together and pushes negatives apart using a similarity score. With cosine similarity $\text{sim}(\mathbf{z}_i, \mathbf{z}_j) = \frac{\mathbf{z}_i \cdot \mathbf{z}_j}{\|\mathbf{z}_i\| \|\mathbf{z}_j\|}$ and temperature τ , define $g(\mathbf{z}_i, \mathbf{z}_j) = \exp(\text{sim}(\mathbf{z}_i, \mathbf{z}_j)/\tau)$. The InfoNCE loss per sample is

$$\text{InfoNCE}(\mathbf{x}_i, \mathbf{x}_i^+) = -\log \left(\frac{g(\mathbf{z}_i, \mathbf{z}_i^+)}{\sum_{j=1}^B g(\mathbf{z}_i, \mathbf{z}_j)} \right). \quad (4)$$

cMIM without Data Augmentation

Intuition. We add a latent Bernoulli k that judges whether $(\mathbf{x}, \mathbf{z})$ is a matched pair. Its calibrated probability yields an augmentation-free contrastive term. We extend the MIM graphical model with k and define the joint factorizations

$$q_\theta(\mathbf{x}, \mathbf{z}, k) = q_\theta(k | \mathbf{x}, \mathbf{z}) q_\theta(\mathbf{z} | \mathbf{x}) q_\theta(\mathbf{x}), \quad p_\theta(\mathbf{x}, \mathbf{z}, k) = p_\theta(k | \mathbf{x}, \mathbf{z}) p_\theta(\mathbf{x} | \mathbf{z}) p_\theta(\mathbf{z}). \quad (5)$$

With $\mathbf{z}_i \sim q_\theta(\mathbf{z} | \mathbf{x}_i)$, the discriminator over k shares parameters in both paths and is Bernoulli with success probability

$$p_{k=1}(\mathbf{x}_i, \mathbf{z}_i) = \frac{g(\mathbf{z}_i, \mathbf{z}_i)}{g(\mathbf{z}_i, \mathbf{z}_i) + \mathbb{E}_{\mathbf{x}' \sim \mathcal{P}(\mathbf{x}), \mathbf{z}' \sim q_\theta(\mathbf{z} | \mathbf{x}')} [g(\mathbf{z}_i, \mathbf{z}')] } \approx \frac{g(\mathbf{z}_i, \mathbf{z}_i)}{g(\mathbf{z}_i, \mathbf{z}_i) + \frac{1}{B-1} \sum_{j=1, j \neq i}^B g(\mathbf{z}_i, \mathbf{z}_j)}. \quad (6)$$

During training, $k=1$ because $\mathbf{z}_i$ is sampled given $\mathbf{x}_i$; the expectation is approximated in-batch.

Concentration. *Intuition.* The in-batch estimate of the negative mean is well-behaved for moderate B . Since cosine similarity lies in $[-1, 1]$, $g \in [e^{-1/\tau}, e^{1/\tau}]$. Hoeffding's inequality implies the in-batch Monte Carlo estimate of the negative mean concentrates around its expectation:

$$\Pr \left(\left| \frac{1}{B-1} \sum_{j \neq i} g(\mathbf{z}_i, \mathbf{z}_j) - \mu \right| \geq \epsilon \right) \leq 2 \exp \left(-\frac{2(B-1)\epsilon^2}{(e^{1/\tau} - e^{-1/\tau})^2} \right). \quad (7)$$

Relation to InfoNCE

Intuition. $-\log p_{k=1}$ is an InfoNCE-like cross-entropy with a *calibrated* positive logit that removes batch-size dependence at logit parity. Let $s_{ij} = \text{sim}(\mathbf{z}_i, \mathbf{z}_j)/\tau$ so $g(\mathbf{z}_i, \mathbf{z}_j) = \exp(s_{ij})$. Then from Eq. (6),

$$p_{k=1} = \frac{\exp(s_{ii})}{\exp(s_{ii}) + \frac{1}{B-1} \sum_{j \neq i} \exp(s_{ij})} = \frac{\exp(s_{ii} + \log(B-1))}{\exp(s_{ii} + \log(B-1)) + \sum_{j \neq i} \exp(s_{ij})}. \quad (8)$$

Thus $-\log p_{k=1}$ equals an InfoNCE cross-entropy where the positive logit is shifted by $\log(B-1)$.

Calibration. If all logits are equal, $p_{k=1} = 1/2$ (independent of B) versus $1/B$ in standard InfoNCE. With cosine similarity $s_{ii} = 1/\tau$ is constant; attraction comes from the MIM term.

211 cMIM Training Objective

212 *Intuition.* We optimize an A-MIM upper bound with an added contrastive term, preserving recon-
 213 struction while shaping global angular geometry. Define the mixture model

$$\mathcal{M}_\theta(\mathbf{x}, \mathbf{z}, \mathbf{k}) = \frac{1}{2} (p_\theta(\mathbf{k} | \mathbf{z}, \mathbf{x}) p_\theta(\mathbf{x} | \mathbf{z}) p_\theta(\mathbf{z}) + q_\theta(\mathbf{k} | \mathbf{z}, \mathbf{x}) q_\theta(\mathbf{z} | \mathbf{x}) q_\theta(\mathbf{x})), \quad (9)$$

214 with sampling distribution $\mathcal{M}_S(\mathbf{x}, \mathbf{z}, \mathbf{k})$ as in MIM . The learning objective upper-bounds the negative
 215 mixture entropy:

$$\mathcal{L}_{\text{MIM}}(\theta) = \frac{1}{2} (\text{CE}(\mathcal{M}_S, q_\theta) + \text{CE}(\mathcal{M}_S, p_\theta)) \geq H_{\mathcal{M}_S}(\mathbf{x}, \mathbf{k}) + H_{\mathcal{M}_S}(\mathbf{z}) - I_{\mathcal{M}_S}(\mathbf{x}, \mathbf{k}; \mathbf{z}). \quad (10)$$

216 For A-MIM (sampling along the encoding path),

$$\begin{aligned} \mathcal{L}_{\text{A-MIM}}(\theta) = -\frac{1}{2} \mathbb{E}_{\mathbf{x} \sim \mathcal{P}(\mathbf{x}), \mathbf{z} \sim q_\theta(\mathbf{z} | \mathbf{x}), \mathbf{k}=1} & \left[\log p_\theta(\mathbf{k} | \mathbf{z}, \mathbf{x}) + \log p_\theta(\mathbf{x} | \mathbf{z}) + \log p_\theta(\mathbf{z}) \right. \\ & \left. + \log q_\theta(\mathbf{k} | \mathbf{z}, \mathbf{x}) + \log q_\theta(\mathbf{z} | \mathbf{x}) + \log q_\theta(\mathbf{x}) \right]. \end{aligned} \quad (11)$$

217 The empirical loss with anchor prior $p(\mathbf{z}) = \mathcal{N}(0, \mathbf{I})$ is

$$\hat{\mathcal{L}}_{\text{A-MIM}}(\theta; \mathcal{D}) = -\frac{1}{N} \sum_{i=1}^N \left(\log p_\theta(\mathbf{x}_i | \mathbf{z}_i) + \log p_{k=1}(\mathbf{x}_i, \mathbf{z}_i) + \frac{1}{2} (\log q_\theta(\mathbf{z}_i | \mathbf{x}_i) + \log p(\mathbf{z}_i)) \right). \quad (12)$$

218 B Experiment and Training Details

219 To evaluate cMIM, we conduct experiments on a 2D toy example, MNIST-like images, and molecular
 220 property prediction (MolMIM by Reidenbach et al. [2023]). The toy example isolates the geometric
 221 effect of the contrastive term. For images, we examine classification, batch-size sensitivity, and
 222 reconstruction. For molecules, we assess reconstruction and downstream regression.

223 B.1 Experiment Details and Datasets

#	Dataset	Train Samples	Test Samples	Categories	Description
1	MNIST	60,000	10,000	10	Handwritten digits
2	Fashion MNIST	60,000	10,000	10	Clothing images
3	EMNIST Letters	88,800	14,800	27	Handwritten letters
4	EMNIST Digits	240,000	40,000	10	Handwritten digits
5	PathMNIST	89,996	7,180	9	Colon tissue histology
6	DermaMNIST	7,007	2,003	7	Skin lesion images
7	OCTMNIST	97,477	8,646	4	Retinal OCT images
8	PneumoniaMNIST	9,728	2,433	2	Pneumonia chest X-rays
9	RetinaMNIST	1,600	400	5	Retinal fundus images
10	BreastMNIST	7,000	2,000	2	Breast tumor ultrasound
11	BloodMNIST	11,959	3,432	8	Blood cell microscopy
12	TissueMNIST	165,466	47,711	8	Kidney tissue cells
13	OrganAMNIST	34,581	8,336	11	Abdominal organ CT scans
14	OrganCMNIST	13,000	3,239	11	Organ CT, central slices
15	OrganSMNIST	23,000	5,749	11	Organ CT, sagittal slices

Table 2: **Image classification datasets.** MNIST/EMNIST (rows 1–4) are handwriting; rows 5–15 are MedMNIST biomedical imaging tasks spanning pathology, retina, chest X-ray, ultrasound, and CT.

224 All models are trained unsupervised. We select the checkpoint with lowest validation loss (no
 225 peeking at test accuracy). For downstream tasks, we freeze the encoder–decoder and train lightweight
 226 classifiers on learned representations. Accuracy is not monitored during pretraining to avoid selection
 227 bias; training runs to convergence for fairness across models.

228 **2D Toy Example.** A synthetic dataset of 1000 points in 2D is initialized in the first quadrant. We
229 visualize how the contrastive term in Eq. (6) shapes latent geometry.

230 **Image Classification on MNIST-like Datasets.** We train MIM, cMIM, VAE, cVAE (VAE +
231 contrastive term), and InfoNCE to convergence on MNIST Deng [2012], EMNIST (letters, digits)
232 Cohen et al. [2017], and MedMNIST Yang et al. [2021]. Images are resized to 28×28 and binarized
233 when needed. We use $\tau = 0.1$ (as in InfoNCE) after a small sweep $\tau \in \{0.1, 1\}$. Encoder:
234 Perceiver Jaegle et al. [2021] with 1 cross-attention, 4 self-attention layers, hidden size 16; 784
235 pixels $\rightarrow$ 400 steps $\rightarrow$ 64-dim latent. Decoder mirrors the encoder. Training: 500k steps; batch-sizes
236 $\{2, 5, 10, 100, 200\}$; Adam lr 10^{-3} ; WSD scheduler Hu et al. [2024]. Classifiers: KNN ($k=5$; cosine
237 & Euclidean) and one-hidden-layer MLP (width 400, Adam 10^{-3} , 1000 steps).

238 **Molecular Property Prediction.** We use ZINC-15 Sterling and Irwin [2015] with SMILES
239 Weininger [1988], following Reidenbach et al. [2023]. Targets: ESOL, FreeSolv, Lipophilicity. Here
240 $\tau = 1$; both MIM and cMIM train for 250k steps. We evaluate SVM/MLP regressors with/without
241 informative embeddings and compare to CDDD Winter et al. [2019]. Additional architectural details
242 are below.

243 B.2 Image Classification

244 Architecture.

- 245 • Encoder: flatten $\rightarrow$ linear to (784, 16) $\rightarrow$ Perceiver to (400, 16) $\rightarrow$ linear to 1 $\rightarrow$ layer norm
246 $\rightarrow$ linear to 64.
- 247 • $q_\theta(z | x)$: Gaussian with mean/variance from linear heads on encoder output.
- 248 • Decoder: linear 64 $\rightarrow$ (64, 16) $\rightarrow$ Perceiver (400, 16) $\rightarrow$ linear to 1 $\rightarrow$ layer norm $\rightarrow$ linear
249 to 784 $\rightarrow$ reshape to 28×28 .
- 250 • $p_\theta(x | z)$: Bernoulli with logits predicted from decoder output.
- 251 • Prior: standard Normal.

252 All models use Adam lr $1e-3$ with WSD (10% warmup/decay) for 500k steps, independent of
253 batch-size.

254 B.3 Molecular Property Prediction

255 **Dataset.** We train on a tranche of ZINC-15 [Sterling and Irwin, 2015], reactive & annotated, with
256 molecular weight ≤ 500 Da and $\log P \leq 5$. We select 730M molecules and split into train/val/test
257 (723M train). To isolate framework effects, we hold architecture and hyperparameters fixed across
258 models. For context, Chemformer used 100M molecules [Sterling and Irwin, 2015] (about $20 \times$
259 CDDD’s 72M from ZINC-15+PubChem [Kim et al., 2018]); MolFormer-XL used ~ 1.1 B molecules.

260 **Data augmentation.** Following Irwin et al. [2022], we use masking and SMILES enu-
261 meration [Weininger, 1988]. Masking (10%) is used only for MegaMolBART. MegaMol-
262 BART/PerBART/MolVAE apply SMILES enumeration with different encoder/decoder permutations;
263 MolMIM improves when encoder and decoder see the *same* permutation, simplifying training.

264 **Model details.** Implemented with NeMo Megatron [Kuchaiev et al., 2019]. RegEx tokenizer with
265 523 tokens [Bird et al., 2009]. Encoders/decoders: 6 layers, hidden size 512, 8 heads, FFN 2048.
266 Perceiver-based models define K (hidden length) with $H = K \times D$ total hidden size (cf. Fig. 2a).
267 Params: MegaMolBART 58.9M; PerBART 64.6M; MolVAE/MolMIM 65.2M. MolVAE uses β -VAE
268 [Higgins et al., 2017] with $\beta = 1/D$.

269 **Optimization.** Adam [Kingma and Ba, 2015] with learning rate 1.0, betas (0.9, 0.999), $\epsilon = 10^{-8}$,
270 weight decay 0. Noam scheduler [Vaswani et al., 2017] (warm-up ratio 0.008; min lr $1e-5$). Max
271 sequence length 512; dropout 0.1; local batch 256; global batch 16384. Training: 1,000,000 steps,
272 fp16, $4 \times$ nodes, $16 \times$ V100 32GB/node. MolVAE also reported with $\beta = 1/H$; we found this
273 balances rate/distortion [Alemi et al., 2018]. MolMIM does not require β tuning.

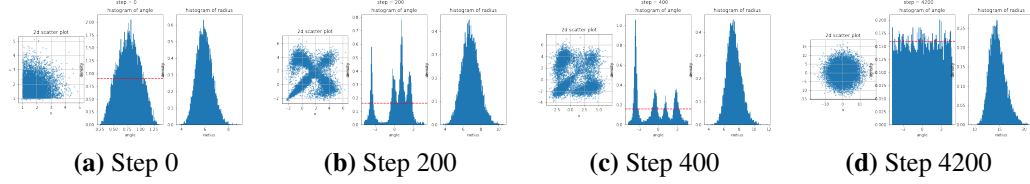


Figure 5: **Contrastive term induces angular uniformity.** Effect of Eq. (6) on a 2D toy example. Each panel shows latent space (left), angle histogram (middle), and radius histogram (right). From (a) initialization to (d) 4200 steps, cMIM distributes points uniformly in angle while allowing radii to vary, complementing MIM’s clustering and improving separability.

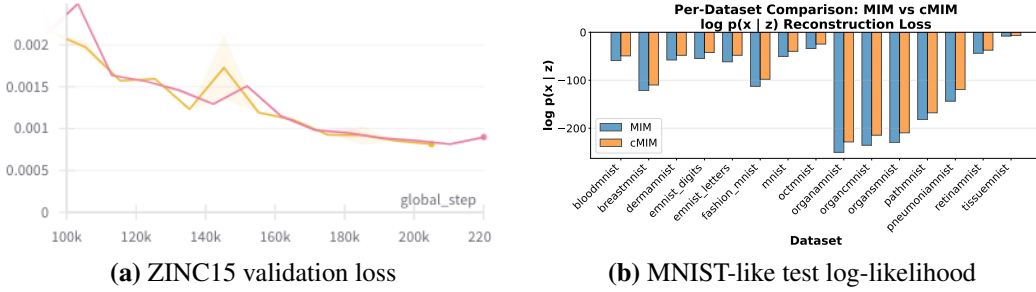


Figure 6: **Generative fidelity is preserved or improved.** (a) Molecular data: cMIM (yellow) and MIM (pink) exhibit comparable validation reconstruction loss. (b) Images: cMIM achieves better average per-example Bernoulli log-likelihood (-96.25) than MIM (-109.64), a 12.2% relative improvement, suggesting a benign regularization effect. (All likelihoods are averaged per example; for images we report mean Bernoulli log-likelihood over pixels.)

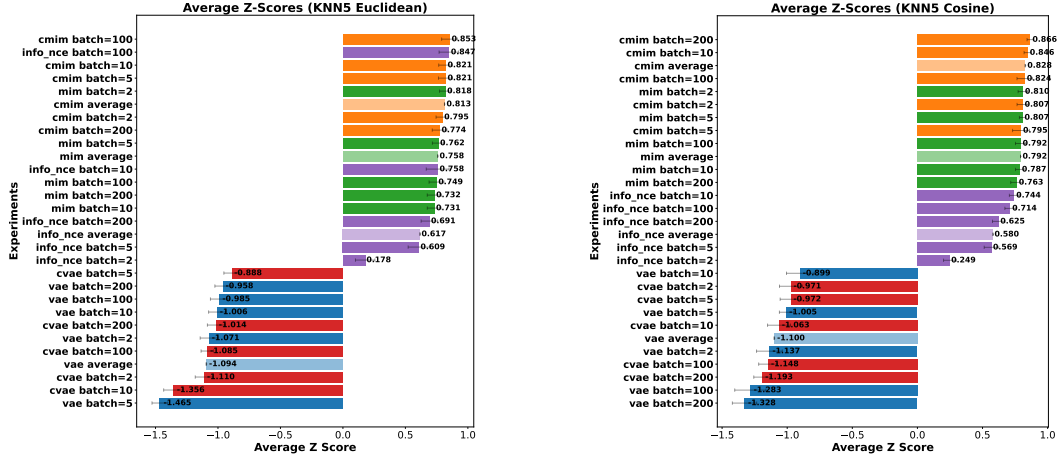
C Additional Results

C.1 Effects of cMIM Loss on 2D Toy Example

We minimize the negative log-likelihood associated with Eq. (6) using $\tau=1$. As expected Wang and Isola [2020], angular uniformity emerges without collapsing radii, indicating that the contrastive term integrates smoothly with the MIM objective.

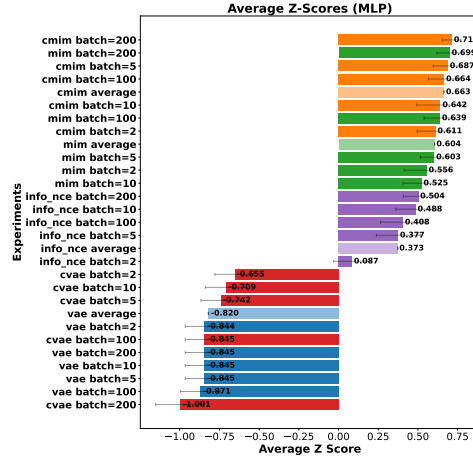
C.2 Reconstruction

C.3 MNIST-like Image Classification



(a) KNN5 (Euclidean)

(b) KNN5 (Cosine)



(c) MLP

Figure 7: **Z-scores across evaluators and batch-sizes.** cMIM attains higher average z -scores than MIM/InfoNCE/VAE across KNN and MLP evaluators, with reduced variance across batch-sizes.

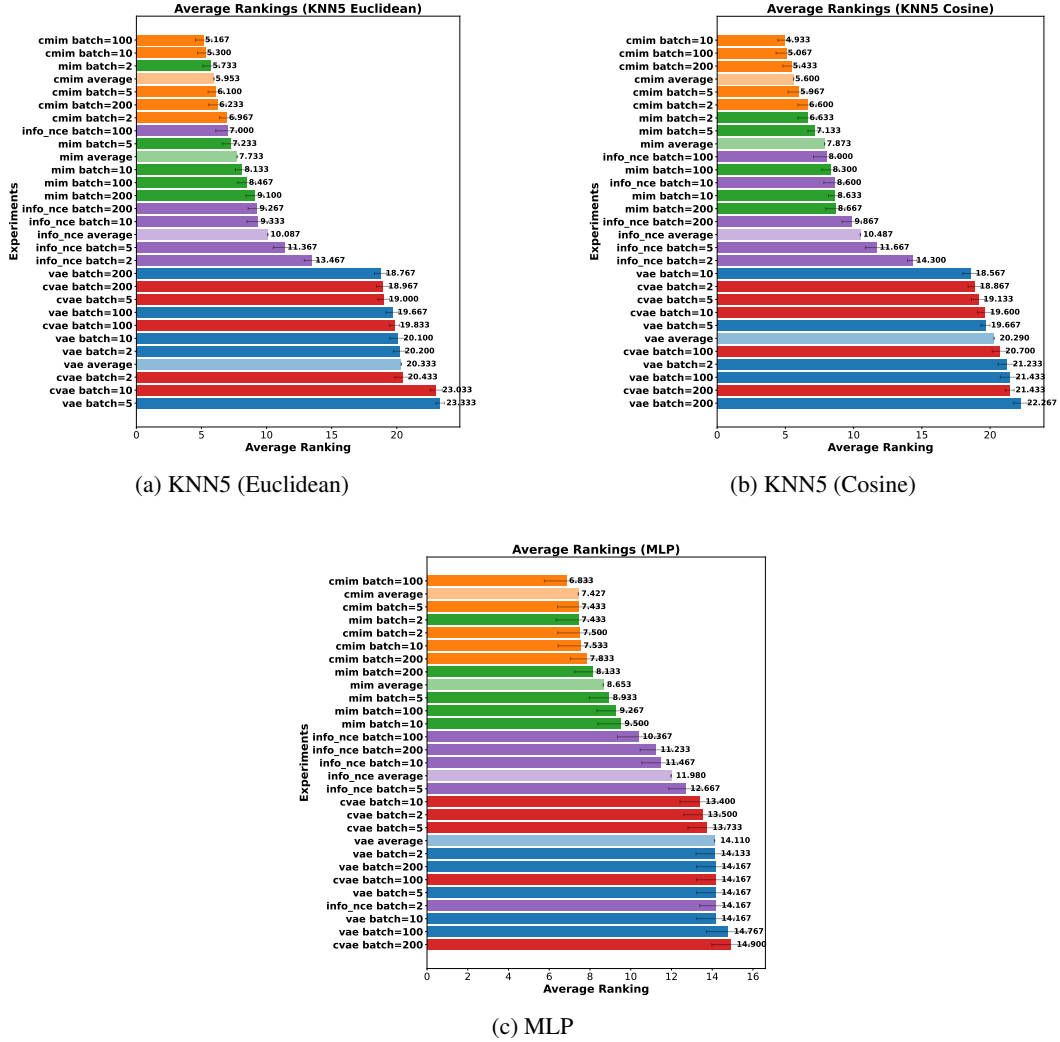


Figure 8: **Average rankings (lower is better).** Across evaluators and batch-sizes, cMIM consistently attains top rankings with narrow error bars, while InfoNCE varies markedly with batch-size.