

i-MAE: ARE LATENT REPRESENTATIONS IN MASKED AUTOENCODERS LINEARLY SEPARABLE?

Anonymous authors

Paper under double-blind review

ABSTRACT

Masked image modeling (MIM) has been recognized as a strong and popular self-supervised pre-training approach in the vision domain. However, the interpretability of the mechanism and properties in the learned representations by such a scheme is so far not well explored. In this work, through comprehensive experiments and empirical studies on Masked Autoencoders (MAE), we address two critical questions to explore the behaviors of the learned representations: **(i)** Are the latent representations in Masked Autoencoders *linearly separable* if the input is a mixture of two images instead of one? This can be concrete evidence to explain why MAE-learned representations have superior performance on downstream tasks, as proven by many literatures impressively. **(ii)** What is the *degree of semantics* encoded in the latent feature space by Masked Autoencoders? To explore these two problems, we propose a simple yet effective *Interpretable MAE* (**i-MAE**) framework with a two-way image reconstruction and a latent feature reconstruction with distillation loss, to help us understand the behaviors inside MAE structure. Extensive experiments are conducted on CIFAR-10/100, Tiny-ImageNet and ImageNet-1K datasets to verify the observations we discovered. Furthermore, in addition to qualitatively analyzing the characteristics in the latent representations, we also examine the existence of linear separability and the degree of semantics in the latent space by proposing two novel metrics. The surprising and consistent results between the qualitative and quantitative experiments demonstrate that i-MAE is a superior framework design for interpretability research of MAE frameworks, as well as achieving better representational ability.

1 INTRODUCTION

Self-supervised learning aims to learn representations from abundant unlabeled data for benefiting various downstream tasks. Recently, many self-supervised approaches have been proposed in the vision domain, such as pre-text based methods (Doersch et al., 2015; Zhang et al., 2016; Gidaris et al., 2018), contrastive learning with Siamese networks (Oord et al., 2018; He et al., 2020; Chen et al., 2020; Henaff, 2020), masked image modeling (MIM) (He et al., 2022; Bao et al., 2022; Xie et al., 2022), etc. Among them, the MIM has shown a preponderant advantage in performance and the representative method Masked Autoencoders (MAE) (He et al., 2022) has attracted much attention in the field. A natural question is raised: *Where is the benefit of the transferability to downstream tasks from in MAE-based training?* This motivates us to develop a framework to shed light on the reasons for the superior latent representation from MAE. Also, as the interpretability of MAE framework is still under-studied in this area, it is crucial to explore this in a specific and exhaustive way.

Intuitively, a good representation should be separable and contain enough semantics from input, so that it can have a qualified ability to distinguish different classes with better performance on downstream tasks. While, how to evaluate the separability and the degree of semantics on latent features is not clear so far. Moreover, the mechanism of an Autoencoder to *compress* the information from input by reconstructing itself, has been a strong self-supervised learning architecture, but the explanation of the learned features through this way is still under-explored.

To address the difficulties of identifying separability and semantics in the latent features, we first propose a novel framework i-MAE upon vanilla MAE. It consists of a mixture-based masked au-

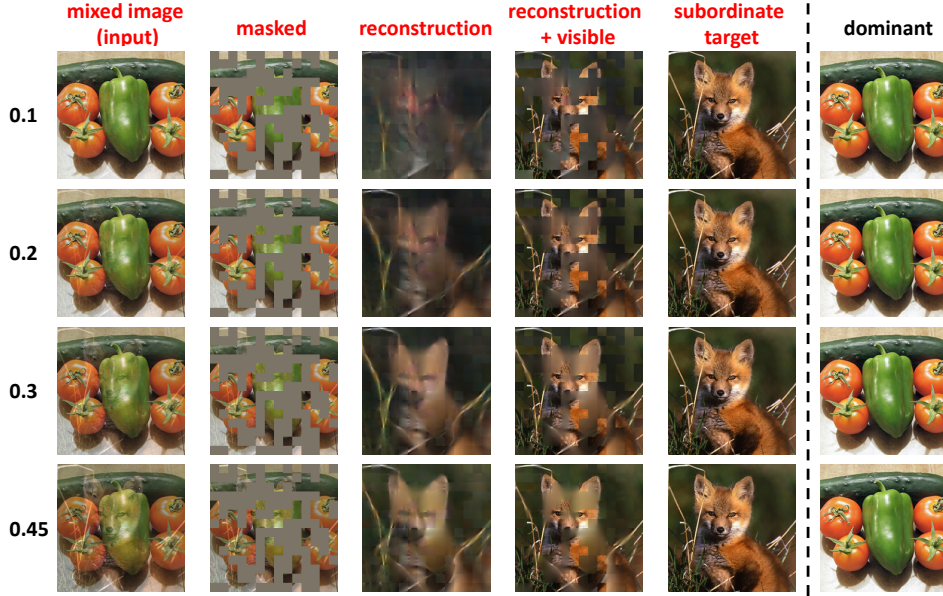


Figure 1: Reconstruction results of *i*-MAE on ImageNet-1K validation images with different mixing coefficients α (listed on the left). *i*-MAE is pre-trained with the subordinate image I_s as the only target and a 0.5 mask ratio. Visually it reflects features of I_s even at 0.1 and still reconstructs the individual image well, whereas at 0.45 reconstructions show the appearance of dominant image (hence the green patches). The input is linearly mixed. More visualizations are provided in Appendix.

toencoder branch for disentangling the mixed representations by linearly separating two different instances, and a pre-trained vanilla MAE as the guidance to distill the disentangled representations. An illustration of the overview framework architecture is shown in Fig. 2. This framework is designed for answering two interesting questions: **(i)** Are the latent representations in Masked Autoencoders *linearly separable*? **(ii)** What is the *degree of semantics* encoded in the latent feature space by Masked Autoencoders? These two questions can reveal the factor that MAE learned features are good at separating different classes. We attribute the superior representation of MAE to it learning separable features for downstream tasks with enough semantics.

In addition to qualitative studies, we also develop two metrics to address the two questions quantitatively. In the first metric, we employ ℓ_2 distance from the high-dimensional Euclidean spaces to measure the similarity between *i*-MAE’s disentangled feature and “ground-truth” feature from pre-trained MAE on the same image. In the second metric, we control different ratios of semantic classes as a mixture within a mini-batch and evaluate the finetuning and linear probing results of the model to reflect the learned semantic information. More details will be provided in Section 3.

We conduct extensive experiments on different scales of datasets: small CIFAR-10/100, medium Tiny-ImageNet and large ImageNet-1K to verify the linear separability and the degree of semantics in the latent representations. We also provide both qualitative and quantitative results to explain our observations and discoveries. The characteristics we observed in latent representations according to our proposed *i*-MAE framework: **(I)** *i*-MAE learned feature representation has good linear separability for input data, which is beneficial for downstream tasks. **(II)** Though the training scheme of MAE is different from instance classification pre-text in contrastive learning, its representation still encodes sufficient semantic information from input data. Moreover, *mixing the same class images as the input substantially improves the quality of learned features*. **(III)** We can reconstruct an image from a mixture by *i*-MAE effortlessly. To the best of our knowledge, this is the pioneer study to explicitly explore the separability and semantics inside MAE’s features with extensive well-designed qualitative and quantitative experiments.

Our contributions in this work are:

- We propose an *i*-MAE framework with two-way image reconstruction and latent feature reconstruction by a distillation loss, to explore the interpretability of mechanisms and properties inside the learned representations of MAE framework.

- We introduce two metrics to examine the linear separability and the degree of semantics quantitatively on the learned latent representations.
- We conduct extensive experiments on different scales of datasets: CIFAR-10/100, Tiny-ImageNet and ImageNet-1K and provide sufficient qualitative and quantitative results.

2 RELATED WORK

Masked image modeling. Motivated by masked language modeling’s success in language tasks (Devlin et al., 2018; Radford & Narasimhan, 2018), Masked Image Modeling (MIM) in the vision domain learn representations from images corrupted by masking. State-of-the-art results on downstream tasks are achieved by several approaches. BEiT (Bao et al., 2022) proposes to recover discrete visual tokens, whereas SimMIM (Xie et al., 2022) addresses the MIM task as a pixel-level reconstruction. In this work, we focus on MAE (He et al., 2022), which proposes to use a high masking ratio and a non-arbitrary ViT decoder. Despite the great popularity of MIM approaches and their conceptual similarity to language modeling, the question of why has not been addressed in the visual domain. Moreover, as revealed by MAE, pixels are semantically sparse, and we novelly examine semantic-level information quantitatively.

Image mixtures. Widely adopted mixture methods in visual supervised learning include Mixup (Zhang et al., 2017) and Cutmix (Yun et al., 2019). However, these methods require ground-truth labels for calculating mixed labels; in this work, we adapt Mixup to our unsupervised framework by formulating losses on only one of the two input images. On the other hand, in very recent visual SSL, joint embedding methods and contrastive learning approaches such as MoCo (He et al., 2020), SimCLR (Chen et al., 2020), and more recently UnMix (Shen et al., 2022) have acquired success and predominance in mixing visual inputs. These approaches promote instance discrimination by aligning features of augmented views of the same image. However, unlike joint embedding methods, i-MAE does not heavily rely on data augmentation and negative sampling. Moreover, whereas most MIM methods are generative tasks, i-MAE also utilizes characteristics of discriminative tasks in learning linearly separable representations.

Invariance and disentangling representation learning in Autoencoders. Representation learning focuses on the properties of the features learned by the layers of deep models while remaining agnostic to the particular optimization process. Variance and entanglement are two commonly discussed factors that occur in data distribution for representation learning. In this work, we focus on the latent disentanglement that one feature is correlated or connected to other vectors in the latent space. Autoencoder is a classical generative unsupervised representation learning framework based on image reconstruction as loss function. Specifically, autoencoders learn both the mapping of inputs to latent features and then the reconstruction of the original input. Denoising autoencoders reconstruct the original input from a corrupted input, and most MIM methods are categorized as denoising autoencoders that use masking as a noise type. We notice, that recent work in the literature (He et al., 2022; Bao et al., 2022) performs many experiments in masking strategies, but to the best of our knowledge, we are the first to introduce image mixtures in the pre-training of MIM.

3 I-MAE

In this section, we first introduce an overview of our proposed framework. Then, we present each component in detail. Ensuing, we elaborate on the metrics we proposed evaluating linear separability and degree of semantics, as well as broadly discuss the observations and discoveries.

3.1 FRAMEWORK OVERVIEW

As shown in Fig 2, our framework consists of three submodules: **(i)** a mixture encoder module that takes the masked mixture image as the input and output mixed features; **(ii)** a disentanglement module that splits the mixed feature to the individual ones; **(iii)** MAE teacher module that provides the pre-trained embedding for guiding the splitting process in the disentanglement module.

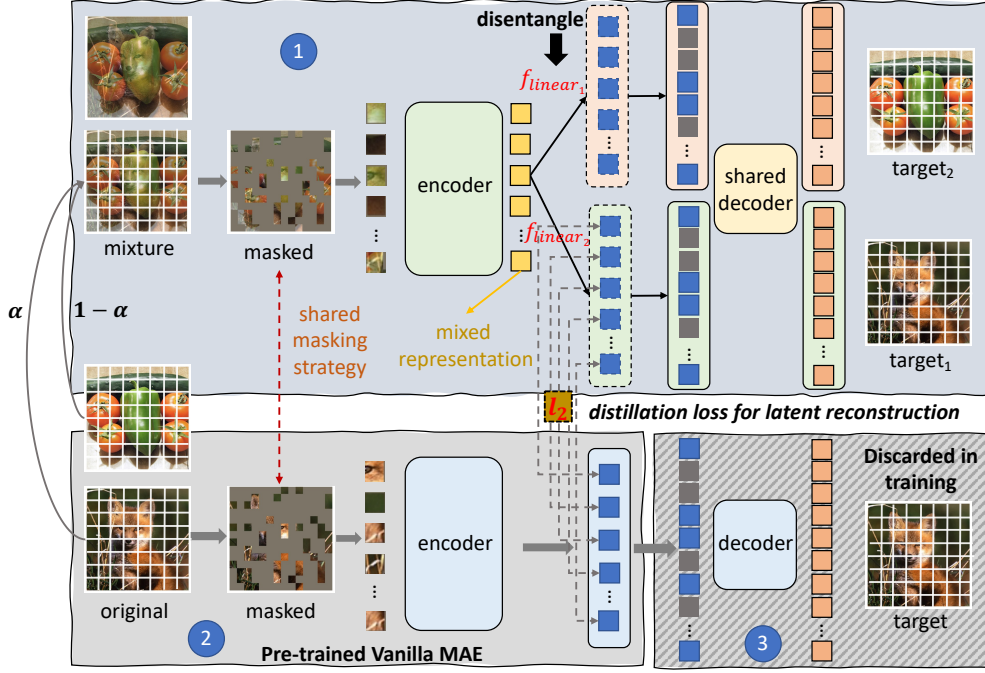


Figure 2: Framework overview of our *i*-MAE. ① is the main branch that consists of a mixture encoder, a disentanglement module, and a two-way image reconstruction module. ② is the encoder part of a pre-trained vanilla MAE for distillation purposes (i.e., latent reconstruction). ③ is the decoder part in MAE and is discarded in training.

3.1.1 COMPONENTS

Input Mixture with MAE Encoder. Inspired by Mixup, we use an unsupervised mixture of inputs formulated by $\alpha * I_1$ and $(1 - \alpha) * I_2$, I_1, I_2 are the input mixes. Essentially, our encoder extrapolates mixed features from a tiny fraction (e.g., 25%) of visible patches, which we then tune to only represent the subordinate image. The mixed image will be:

$$I_m = \alpha * I_1 + (1 - \alpha) * I_2 \quad (1)$$

where α is the coefficient to mix two images following a Beta distribution.

Two-branch Masked Autoencoders with Shared Decoder. Although sufficient semantic information of both images is embedded in the mixed representation to reconstruct both images, the vanilla MAE cannot by itself associate separated features to either input. The MAE structure does not retain identification information of the two mixed inputs (e.g., order or positional information), i.e., the model cannot tell which of the two images to deconstruct to, since both are sampled from the same distribution and mixed randomly. The consequence is that both reconstructions look identical to each other and fail to look similar to either original input.

Similar to how positional embeddings are needed to explicitly encode spatial information, *i*-MAE implicitly encodes the semantic difference between the two inputs by using a dominant and subordinate mixture strategy. Concretely, through an unbalanced mix ratio and a reconstruction loss targeting only one of the inputs, our framework encodes sufficient information for *i*-MAE to linearly map the input mixture to two outputs.

Two-way Image Reconstruction Loss. Formally, we build our reconstruction loss to recover individual images from a mixed input, which is first fed into the encoder to generate mixed features:

$$h_m = E_{i\text{-MAE}}(I_m) \quad (2)$$

where $E_{i\text{-MAE}}$ is *i*-MAE’s encoder, h_m is the latent mixed representation. Then, we employ two non-shareable linear embedding layers to separate the mixed representation from the individual ones:

$$\begin{aligned} h_1 &= f_1(h_m) \\ h_2 &= f_2(h_m) \end{aligned} \quad (3)$$

where f_1, f_2 are two linear layers with different parameters for disentanglement and h_1 and h_2 are corresponding representations. After that, we feed the individual representations into the shared decoder with the corresponding reconstruction losses:

$$\begin{aligned}\mathcal{L}_{\text{recon}}^{I_1} &= \mathbb{E}_{I_1 \sim p(I_1)} [\|D_{\text{shared}}(h_1) - I_1\|_2] \\ \mathcal{L}_{\text{recon}}^{I_2} &= \mathbb{E}_{I_2 \sim p(I_2)} [\|D_{\text{shared}}(h_2) - I_2\|_2]\end{aligned}\quad (4)$$

In practice, we train the linear separation layers to distinguish between the dominant input I_d (higher mix ratio) and the subordinate input I_s (lower ratio). Showing that our encoder learns to embed representations of both images, we intentionally choose to reconstruct only the subordinate image I_s to prevent the I_d from guiding the reconstruction. Essentially, successful reconstructions from only the I_s prove that representations of both images can be learned and that the subordinate image is not filtered out as noise.

Patch-wise Distillation Loss for Latent Reconstruction. With the linear separation layers and an in-balanced mixture, the i-MAE encoder is presented with sufficient information about both images to perform visual reconstructions; however, information is inevitably lost during the mixture process, harming the value of the learned features in downstream tasks such as classification. To mitigate such an effect, we propose a knowledge distillation module both for enhancing the learned feature’s quality, and that a successful distillation can evidently prove the linear separability of our features.

Intuitively, MAE’s feature can be regarded as ”ground-truth” and i-MAE learns features distilled from the original MAE. Specifically, our loss function computes ℓ_2 loss between disentangled representations and original representations to help our encoder learn useful features of both inputs.

Our Patch-wise latent reconstruction loss can be formulated as:

$$\begin{aligned}\mathcal{L}_{\text{recon}}^{h_1} &= \mathbb{E}_{h_1 \sim q(h_1)} [\|E_{\text{p-MAE}}(I_1) - h_1\|_2] \\ \mathcal{L}_{\text{recon}}^{h_2} &= \mathbb{E}_{h_2 \sim q(h_2)} [\|E_{\text{p-MAE}}(I_2) - h_2\|_2]\end{aligned}\quad (5)$$

where $E_{\text{p-MAE}}$ is the pre-trained MAE encoder.

3.2 LINEAR SEPARABILITY

For i-MAE to reconstruct the subordinate image from a linear mixture, not only does the encoder have to be general enough to retain information of both inputs, but it must also generate embeddings that are specific enough for the decoder to distinguish them into their pixel-level forms. A straight-forward interpretation of how i-MAE fulfills both conditions is that the latent mixture h_m is a linear combination of features that closely relate to h_1 and h_2 , e.g. in a linear relationship. Our distillation module aids the information loss. To verify this explanation, we employ a linear separability metric to experimentally observe such behavior.

Metric of Linear Separability. A core contribution of our i-MAE is the quantitative analysis of features. In general, linear separability is a property of two sets of features that can be separated into their respective sets by a hyperplane. In our example, the set of latent representations H_1 and H_2 are linearly separable if there exists $n + 1$ real numbers $w_1, w_2, \dots, w_n, b$, such that every $h \in H_1$ satisfies $\sum w_i h_i > b$ and every $h \in H_2$ satisfies $\sum w_i h_i < b$. It is a common practice to train a classical linear classifier (e.g., SVM) and evaluate if two sets of data are linearly separable.

However, to quantitatively measure the separation of latent representations, we devised a more intuitive yet effective metric. Our metric computes the Mean Squared Error (MSE) distance between the disentangled feature of the subordinate image I_s and the vanilla MAE feature of a single input I_s . Since the disentangled feature without constraints will unlikely resemble the vanilla feature, we utilize a linear layer to transform the disentangled feature space to the vanilla feature space. Note that this is similar to knowledge distillation, but happens after the pre-training process without finetuning the parameters and conceptually measures the distance between the two latent representations, and thus the linear transformation will not be needed for i-MAE with distillation. The detailed formulation of the metric is:

$$\mathcal{M}_{\text{ls}} = \frac{1}{N} \sum_{n=1}^N \|h_s^n - f_{\theta}(I_s^n)\|_2^2 \quad (6)$$

where N is the total number of samples. f_{θ} is the encoder of vanilla MAE. I_s is the subordinate image and $I_s \in \{I_1, I_2\}$.

3.3 SEMANTICS

Metric of Semantics. Vanilla MAE exhibits strong signs of semantics understanding (He et al., 2022). However, studying the abstract concept of semantics in the visual domain is difficult due to its semantic sparsity and repetitiveness. Addressing this problem, we propose a metric unique to i-MAE that is readily available for examining the degree of semantics learned in the model. Asides from straightforwardly evaluating classification accuracy to measure the quality of latent representation, i-MAE utilizes the mixing of semantically similar instances to determine to what degree the disentangled latent representations can reflect image-level meaning.

Naturally, the segmentation of different instances from the same class is a more difficult task than classification between different classes, intra-class separation requires the understanding of high-level visual concepts, i.e. semantic differences, rather than lower-level patterns, i.e. shape or color. While generally, data transformation (Olah et al., 2017) can help mitigate overfitting. Similarly, our semantic disentangle module is another data augmentation that introduces significantly more mixtures of the same class into the training process. We find our method to boost the semantics of features learned by this *semantics-controllable mixture* scheme. Specifically, we choose training instances from the same or different classes following different distributions to constitute an input mixture, so that to examine the quality of learned features as follows:

$$\mathbf{p} = \mathbf{f}_m(\mathbf{I}_{c_a} + \mathbf{I}_{c_b}) \quad (7)$$

where $\mathbf{f}_m$ is the backbone network for mixture input and $\mathbf{p}$ is the corresponding prediction. $\mathbf{I}_{c_a}$ and $\mathbf{I}_{c_b}$ are the input samples and c_a, c_b have a certain percentage r that belongs to the same category. For instance, if $r = 0.1$, it indicates that 10% images in a mini-batch are mixed with the same class. When $r = 1.0$, all training images will be mixed by another one from the same class, which can be regarded as a semantically enhanced augmentation. During training, r is fixed for individual models, and we study the degree of semantics that the model encoded by changing the percentage value r . After the model is trained by i-MAE using such kind of input data, we finetune the model with Mixup strategy (both baseline and our models) and cross-entropy loss. We use the accuracy as the metric of semantics under this percentage of instance mixture:

$$\mathcal{M}_{\text{sem}} = - \sum_{i=1}^n t_i \log(\mathbf{p}_i) \quad (8)$$

where t_i is the ground-truth. The insight behind is that: if the input mixture is composed of two images or instances with the same semantics (i.e., the same category), it will confuse the model during training and i-MAE will be struggling to disentangle them. Thus, the encoded information/semantics may be weakened in training and it can be reflected by the quality of learned representation. It is interesting to see whether this conjecture is supported by the empirical results. We use the representation quality through finetuned accuracy to monitor the degree of semantics with this *semantics-controllable mixture* scheme.

4 EMPIRICAL RESULTS AND ANALYSIS

In the experiments section, we analyze the properties of i-MAE’s disentangled representations on an extensive range of datasets. First, we provide the datasets used and our implementations details. Then, we thoroughly ablate our experiments, focusing on the properties of *linear separation*, and *controllable-semantic mixture*. Lastly, we give the final evaluation of our results.

4.1 DATASETS AND TRAINING IMPLEMENTATION FOR BASELINE AND I-MAE.

Settings: We perform empirical experiments of i-MAE on CIFAR-10/100, Tiny-ImageNet, and ImageNet-1K. On CIFAR-10/100, we pre-train i-MAE unsupervisedly and adjust MAE’s structure to better fit the smaller datasets: ViT-Tiny (Touvron et al., 2021) in the encoder and a lite-version of ViT-Tiny (4 layers) as the decoder. Our pre-training lasts 2,000 epochs with learning rate 1.5×10^{-4} and 200 warm-up epochs. On Tiny-ImageNet, i-MAE’s encoder is ViT-small and decoder is ViT-Tiny, trained for 1,000 epochs with learning rate 1.5×10^{-4} . Additionally, we apply warm-up for first 100 epochs, and use cosine learning rate decay with AdamW optimizer as in MAE.

Supervised Finetuning: In the finetuning process, we apply Mixup for all experiments to fit our pre-training scheme, and compare our results with baselines of the same configuration. On CIFAR-10/100, we finetune 100 epochs with a learning rate of 1.5×10^{-3} and AdamW optimizer.

Linear Probing: For linear evaluation, we follow MAE (He et al., 2022) to train with no extra augmentations and use zero weight decay. We also adopt an additional BatchNorm layer without affine transformation.

4.2 ABLATION STUDY

In this section, we perform ablation studies on i-MAE to demonstrate the invariant property of linear separability and to what extent can i-MAE separate features. Then, we analyze the effect of semantic-level instance mixing on the quality of i-MAE’s learned representations.

4.2.1 ABLATION FOR LINEAR SEPARABILITY

To begin, we thoroughly ablate our experiments on small-scaled datasets and demonstrate how i-MAE’s learned features display linear separability. Specifically, we experiment with the separability of the following aspects of our methods: (i) constant or probability mix factor; (ii) masking ratio of input mixtures; (iii) different ViT architectures. Unless otherwise stated, the default settings used in our ablation experiments are ViT-Tiny, masking ratio of 75%, fixed mixing ratio of 35%, and reconstructing only the subordinate image for a harder task.

Mix Ratio. To demonstrate the separable nature of input mixtures, we compared different mixture ratios and random mixture ratios from a Beta distribution. Intuitively, low mixing ratios contain less information that are easily confused with noise, whereas higher ratios destroy the subordinate-dominant relationship. Experimentally, we observe matching results shown in Appendix (Fig. 10 and Fig. 1). The better separation performance around the 0.3 range indicates that i-MAE features are better dichotomized when balanced between noise and useful information. Whereas below 0.15, the subordinate image is noisy and reconstructions are not interpretable, mixing ratios above 0.45 break the balance between the two images, and the two features cannot be distinguished. Moreover, notice that at 0.45, reconstruction patches are turning green and resembling the pepper.

Mask Ratio. In i-MAE, visible information of the subordinate image is inherently limited due to the unbalanced mix ratio in addition to masking. Therefore, a high masking ratio (75% (He et al., 2022)) may not be necessary to suppress the amount of information the encoder sees, so we attempt ratios of 50%, 60% to introduce more information of the subordinate target. As shown in Fig. 3, a lower masking ratio can improve the reconstruction quality.

Combining our findings in mix and mask ratios, we empirically find that i-MAE can compensate for the information loss at low ratios with the additional alleviation of more visible patches (lower mask ratio). Illustrated in Fig. 1, we display a case of i-MAE qualitatively succeeding in separating the features of a $\alpha = 0.1$ mix and 0.5 masking ratio.

Our core finding in the separability ablation section is that i-MAE can learn linearly separable features under two conditions: **(i)** enough information about both images must be present (this can be alleviated by mask ratios). **(ii)** the image-level distinguishing relationship between minority and majority (determined by mix ratio) is clear enough.

ViT Backbone Architecture. We study the effect of different ViT scales in linear separation in Appendix of Fig. 5, and find that larger backbones are not necessary for small datasets on i-MAE, although it is crucial on large-scale ImageNet-1K.

4.2.2 ABLATION FOR DEGREE OF SEMANTICS

Semantic Mixes. Depending on the number of classes and overall size, pristine datasets usually contain around 10% (e.g., CIFAR-10) to less than 1% (e.g. ImageNet-1K) samples of the same class. By default, uniformly random sampling mixtures will be of the same likelihood. However, in the *semantics-controllable mixture* scheme, we test whether the introduction of semantically homogeneous mixtures affects the classification performance at different amounts. We intentionally test to see if similar instances during pre-training can negatively affect the classification performance.

As shown in Tab. 1, after i-MAE pre-training, we perform finetuning and linear probing on classification tasks to evaluate the degree of semantics learned given different amounts of intra-class mix

Table 1: Classification performance of models pre-trained with *semantics-controllable mixture* with intra-class mix rate r from 0.0 to 1.0. Whereas the former are all different pairs, $r = 1$ represents training with mixtures from only the same class.

Same Class Ratio	CIFAR-10		CIFAR-100		Tiny-ImageNet	
	Finetune	Linear	Finetune	Linear	Finetune	Linear
baseline	90.78	72.47	68.66	32.57	59.28	-
0.0	91.67	70.53	-	-	60.91	-
0.5	92.34	72.80	69.50	30.11	60.58	-
1.0	91.60	77.61	69.33	33.39	61.13	-



Figure 3: Comparisons between different mask ratios on Tiny-ImageNet validation dataset. We find i-MAE benefiting from lower masking ratios when reconstructing images.

r . From Tab. 1, we discover that i-MAE overall has a stronger performance in finetuning and linear probing with a non-zero same class ratio. Specifically, a high r increases the accuracy in linear evaluation most in all datasets, meaning that the quality of learned features is best and separated. On the other hand, setting $r = 0.5$ is advantageous during finetuning, as it gains a balanced prior of separating both intra- and inter-class mixtures.

Table 2: Linear separation metric using ℓ_2 distance calculated before and after linear regression on CIFAR-10, CIFAR-100, Tiny-ImageNet, and ImageNet. Reported results are from 100 samples trained for 2000 (5000 on ImageNet-1K) epochs and a fixed mix ratio of 0.3. Baseline is embedding from MAE encoder without layernorm, whereas i-MAE and i-MAE without distillation are embedding after disentanglement.

	CIFAR-10		CIFAR-100		Tiny-ImageNet		ImageNet-1K	
	Before	After	Before	After	Before	After	Before	After
Baseline	3.5899	0.0584	3.0840	0.0487	10.38	0.0204	13.91	0.2004
i-MAE w/o distill	0.1384	0.0568	0.2999	0.0475	0.0723	0.0363	0.8799	0.1316
i-MAE	0.0331	0.0474	0.0312	0.0456	0.0708	0.0352	0.2760	0.1838

4.3 RESULTS OF FINAL EVALUATION

In this section, we provide a summary of our main findings: how separable are i-MAE embedded features and the amount of semantics embedded in mixed-representations. Then, we evaluate the quality of our features with classification and analyze the features.

4.3.1 SEPARABILITY

In this section, we show how i-MAE displays properties of linear separability, visually and quantitatively, and demonstrate our advantage over baseline (vanilla MAE). In a visual comparison of disentanglement capability, shown in Fig. 4, the vanilla MAE does not perform well out-of-the-box. In fact, the reconstructions represent the mixed input more so than the subordinate image.

Since the mixture inputs of i-MAE is a linear combination of the two images and our results show i-MAE’s potent ability to reconstruct both images, even at very low mixture ratios, we account such ability to i-MAE’s disentanglement correlating strongly to vanilla MAE’s features.

As aforementioned, we gave the formal definition of linear separability; we now empirically illustrate the strength of the linear relationship between MAE’s features and i-MAE’s disentangled features with a linear regressor. We employ ℓ_2 distance as our criterion and results are reported in Tab. 2. Experimentally, we feed mixed inputs to i-MAE and singular image to the target model

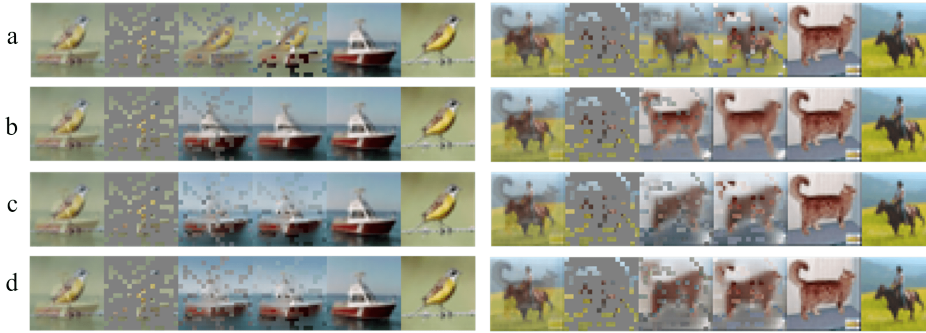


Figure 4: Qualitative comparison on CIFAR-10. **Row (a):** baseline vanilla MAE; **(b):** MAE with unmixed input; **(c):** our i-MAE without distillation; and **(d):** i-MAE with distillation.

Table 3: Finetuning classification acc. on i-MAE with best *semantics-controllable mixture* settings.

Method	CIFAR-10	CIFAR-100	Tiny-ImageNet
Baseline	90.78	68.66	59.28
i-MAE	92.00	69.50	61.63

Table 4: Linear Evaluation accuracy of i-MAE with best *semantics-controllable mixture* settings.

Method	CIFAR-10	CIFAR-100
Baseline	72.47	32.57
i-MAE	77.61	33.39

(vanilla MAE), *Before* indicates that we directly calculate the distance between the “ground-truth” features from pre-trained MAE and our disentangled features. *After* indicates that we train the linear regression’s parameters to fit the “ground-truth”. *Baseline* is the model trained without disentanglement module. It can be observed that our i-MAE has a significantly smaller distance than the vanilla model, reflecting that such a scheme can obtain better separability ability.

4.3.2 SEMANTICS

Finetune and Linear Evaluation. We evaluate our i-MAE’s performance with finetuning and linear evaluation of regular inputs and targets. For all approaches in the finetuning phase, we use Mixup as augmentation and no extra augmentations for linear evaluation. Classification performance is outlined in Tab. 3 and Tab. 4. As our features are learned from a harder scenario, it encodes more information with a more robust representation and classification accuracy. Besides, i-MAE shows a considerable performance boost with both evaluation methods.

Analysis. We emphasize that our enhanced performance comes from i-MAE’s ability to learn more separable features with the disentanglement module, and the enhanced semantics learned from training with *semantics-controllable mixture*. Our classification results show that it is crucial for MAE to learn features that are linearly separable, which can help identify between different classes. However, to correctly identify features with their corresponding classes, semantically rich features are needed, and can be enhanced by sampling intra-class mixing strategy.

5 CONCLUSION

It is non-trivial to understand why Masked Image Modeling (MIM) in the self-supervised scheme can learn useful representations for downstream tasks without labels. In this work, we have introduced a novel interpretable framework upon Masked Autoencoders (i-MAE) to explore two critical properties in latent features: *linear separability* and *degree of semantics*. We identified that the two specialties are the core for superior latent representations and revealed the reasons where is the good transferability of MAE from. Moreover, we proposed two metrics to evaluate these two specialties quantitatively. Extensive experiments are conducted on CIFAR-10/100, Tiny-ImageNet, and ImageNet-1K datasets to demonstrate our discoveries and observations in this work. We also provided sufficient qualitative results and analyses of different hyperparameters. We hope this work can inspire more studies on the interpretability of the MIM frameworks in the future.

REFERENCES

- Hangbo Bao, Li Dong, and Furu Wei. Beit: Bert pre-training of image transformers. In *ICLR*, 2022.
- Ting Chen, Simon Kornblith, Mohammad Norouzi, and Geoffrey Hinton. A simple framework for contrastive learning of visual representations. In *International conference on machine learning*, pp. 1597–1607. PMLR, 2020.
- Jia Deng, Wei Dong, Richard Socher, Li-Jia Li, Kai Li, and Li Fei-Fei. Imagenet: A large-scale hierarchical image database. In *2009 IEEE conference on computer vision and pattern recognition*, pp. 248–255. Ieee, 2009.
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina N. Toutanova. Bert: Pre-training of deep bidirectional transformers for language understanding. 2018. URL <https://arxiv.org/abs/1810.04805>.
- Carl Doersch, Abhinav Gupta, and Alexei A Efros. Unsupervised visual representation learning by context prediction. In *Proceedings of the IEEE international conference on computer vision*, pp. 1422–1430, 2015.
- Spyros Gidaris, Praveer Singh, and Nikos Komodakis. Unsupervised representation learning by predicting image rotations. In *International Conference on Learning Representations*, 2018.
- 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*, pp. 9729–9738, 2020.
- Kaiming He, Xinlei Chen, Saining Xie, Yanghao Li, Piotr Dollár, and Ross Girshick. Masked autoencoders are scalable vision learners. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pp. 16000–16009, 2022.
- Olivier Henaff. Data-efficient image recognition with contrastive predictive coding. In *International conference on machine learning*, pp. 4182–4192. PMLR, 2020.
- Alex Krizhevsky. Learning multiple layers of features from tiny images. 2009.
- Chris Olah, Alexander Mordvintsev, and Ludwig Schubert. Feature visualization. *Distill*, 2017. doi: 10.23915/distill.00007. <https://distill.pub/2017/feature-visualization>.
- Aaron van den Oord, Yazhe Li, and Oriol Vinyals. Representation learning with contrastive predictive coding. *arXiv preprint arXiv:1807.03748*, 2018.
- Adam Paszke, Sam Gross, Francisco Massa, Adam Lerer, James Bradbury, Gregory Chanan, Trevor Killeen, Zeming Lin, Natalia Gimelshein, Luca Antiga, Alban Desmaison, Andreas Kopf, Edward Yang, Zachary DeVito, Martin Raison, Alykhan Tejani, Sasank Chilamkurthy, Benoit Steiner, Lu Fang, Junjie Bai, and Soumith Chintala. Pytorch: An imperative style, high-performance deep learning library. In H. Wallach, H. Larochelle, A. Beygelzimer, F. d’Alché-Buc, E. Fox, and R. Garnett (eds.), *Advances in Neural Information Processing Systems 32*, pp. 8024–8035. Curran Associates, Inc., 2019.
- Alec Radford and Karthik Narasimhan. Improving language understanding by generative pre-training. 2018.
- Zhiqiang Shen, Zechun Liu, Zhuang Liu, Marios Savvides, Trevor Darrell, and Eric Xing. Un-mix: Rethinking image mixtures for unsupervised visual representation learning. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 36, pp. 2216–2224, 2022.
- Hugo Touvron, Matthieu Cord, Matthijs Douze, Francisco Massa, Alexandre Sablayrolles, and Herve Jegou. Training data-efficient image transformers and distillation through attention. In *International Conference on Machine Learning*, volume 139, pp. 10347–10357, July 2021.
- Zhenda Xie, Zheng Zhang, Yue Cao, Yutong Lin, Jianmin Bao, Zhuliang Yao, Qi Dai, and Han Hu. Simmim: A simple framework for masked image modeling. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2022.

- Sangdo Yun, Dongyoon Han, Seong Joon Oh, Sanghyuk Chun, Junsuk Choe, and Youngjoon Yoo. Cutmix: Regularization strategy to train strong classifiers with localizable features. In *Proceedings of the IEEE/CVF international conference on computer vision*, pp. 6023–6032, 2019.
- Hongyi Zhang, Moustapha Cisse, Yann N Dauphin, and David Lopez-Paz. mixup: Beyond empirical risk minimization. *arXiv preprint arXiv:1710.09412*, 2017.
- Richard Zhang, Phillip Isola, and Alexei A Efros. Colorful image colorization. In *European conference on computer vision*, pp. 649–666. Springer, 2016.

APPENDIX

In the appendix, we elaborate on the details and provide more examples of our main text, specifically:

- Section A “Datasets”: specifications of the datasets we used.
- Section B “Implementation Details”: implementation details and configuration settings for unsupervised pre-training and supervised classification.
- Section C “Visualization”: additional reconstruction examples on all datasets.
- Section D “Pseudocode”: a PyTorch-like pseudocode for our mixture and loss methods.

A DATASETS

CIFAR-10/100 (Krizhevsky, 2009) Both CIFAR datasets contain 60,000 tiny colored images sized 32×32 . CIFAR-10 and 100 are split into 10 and 100 classes, respectively.

Tiny-ImageNet The Tiny-ImageNet is a scaled-down version of the standard ImageNet-1K consisting of 100,000 64×64 colored images, categorized into 200 classes.

ImageNet-1K (Deng et al., 2009) The ILSVRC 2012 ImageNet-1K classification dataset consist of 1.28 million training images and 50,000 validation images of 1000 classes.

B IMPLEMENTATION DETAILS IN SELF-SUPERVISED PRE-TRAINING, FINETUNING, AND LINEAR EVALUATION

ViT architecture. In our non-ImageNet datasets, we adopt smaller ViT backbones that generally follow (Touvron et al., 2021). The central implementation of linear separation happens between the MAE encoder and decoder, with a linear projection layer for each branch of reconstruction. A shared decoder is used to reconstruct both images. A qualitative evaluation of different ViT sizes on Tiny-ImageNet is displayed in Fig. 5; the perceptive difference is not large and generally, ViT-small/tiny are sufficient for non-ImageNet datasets.

Pre-training. The default setting for pre-training is listed in Tab. 5. On ImageNet-1K, we strictly use MAE’s specifications. For better classification performance, we use normalized pixels (He et al., 2022) and a high masking ratio (0.75); for better visual reconstructions, we use a lower masking ratio (0.5) without normalizing target pixels. In CIFAR-10/100, and Tiny-ImageNet, reconstruct ordinary pixels.

Semantics-controllable mixture The default setting for our semantics-controllable mixtures are listed in Tab. 6. We modified the dataloader to mix, within a mini-batch, r percent of samples that have homogenous classes, and $1 - r$ percent that is different.

Classification For the classification task, we provide the detailed settings of our finetuning process in Tab. 7 and linear evaluation process in Tab. 8.

C VISUALIZATION

We provide extra examples of a single-branch trained i-MAE reconstructing the subordinate image. Fig. 10 are visualizations on CIFAR-100 at mix ratios from 0.1 to 0.45, in 0.05 steps. Shown in Fig. 6 and Fig. 7, we produce finer ranges of reconstructions from 0.05 to 0.45. Notice that in most cases, mixture rates above 0.4 tends to show features of the dominant image. This observation demonstrates that a low mixture rate can better embed important information separating the subordinate image.

D PYTORCH (PASZKE ET AL., 2019) STYLED PSEUDOCODE

The pseudocode of our mixture and subordinate reconstruction approach is shown in Algorithm 1. This is only a simple demonstration of our most basic framework without distillation losses. In

Algorithm 1: PyTorch-style pseudocode for a single subordinate reconstruction on i-MAE.

```

# alpha: mixture ratio
# args.beta: hyperparameter for the Beta Distribution.
#
# args.beta=1.0
for x in loader: # Minibatch x of N samples
    alpha = np.random.beta(args.beta, args.beta)
    sub_idx = np.argmin(alpha, 1-alpha) # Identifying the
        subordinate (target) image
    perm = torch.randperm(batch_size) # inner-batch mix
    im_1, im_2 = x, x[perm, :]
    mixed_images = alpha * im_1 + (1-alpha) * im_2
#
# Subordinate Loss
loss_sub = loss_fn (model(mixed_images), im_2)
#
# update gradients
optimizer.zero_grad()
loss.backward()
optimizer.step()
...

```

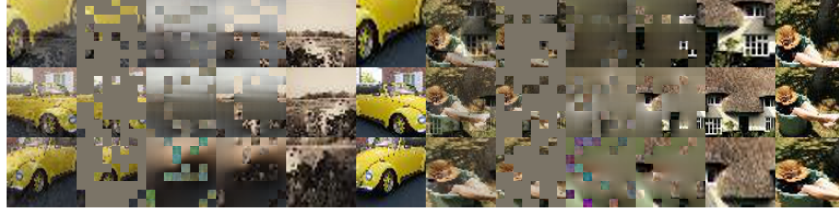


Figure 5: Different ViT backbones (tiny, small, and base) on Tiny-ImageNet . Reconstruction quality is moderately improved when a larger backbone is used.

our full-fledged i-MAE, we employ two additional distillation losses, an additional linear separation branch, and the *semantics-controllable mixture* scheme; nonetheless, the key implementation remains the same as the pseudocode presented here.

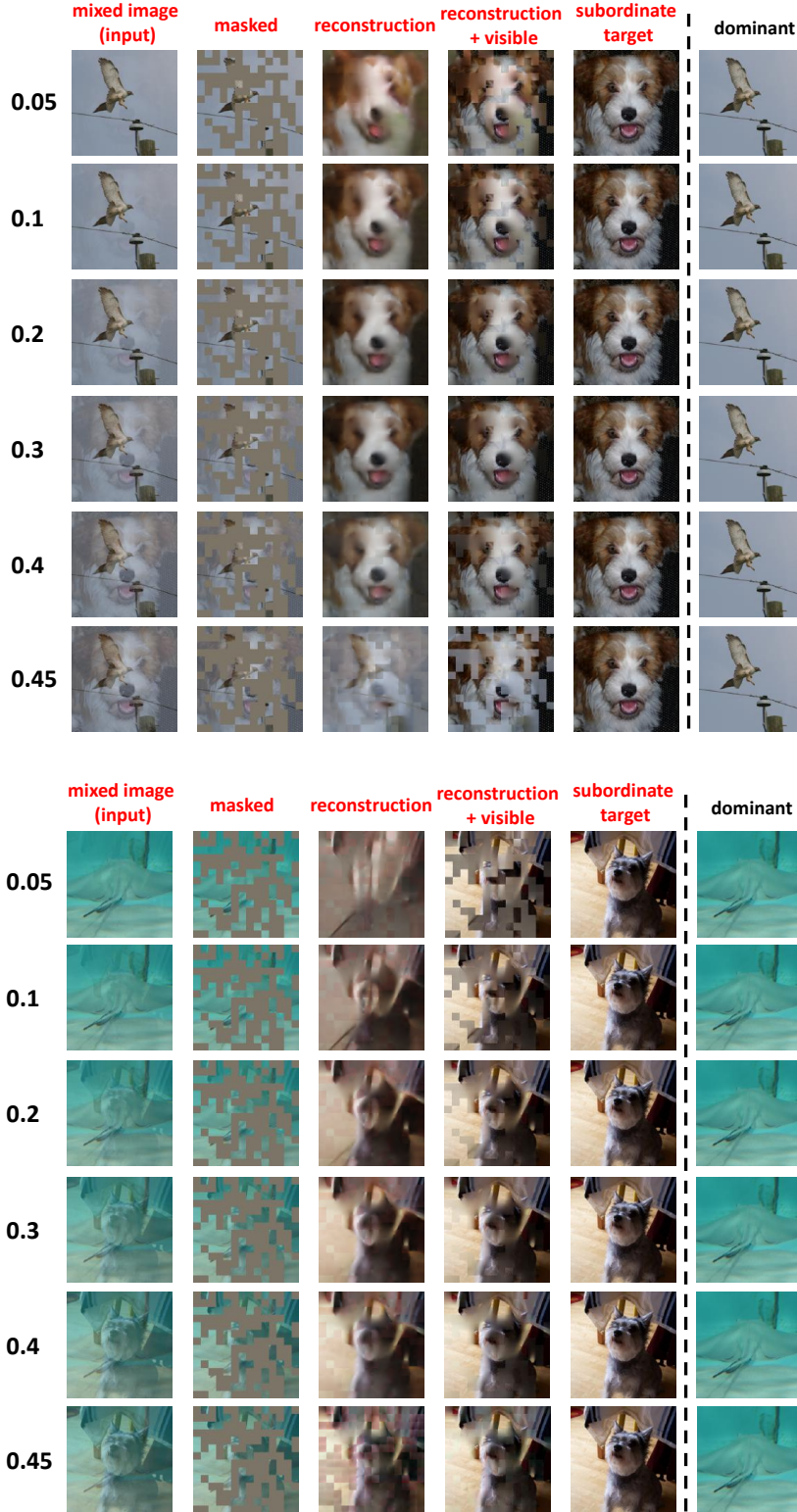


Figure 6: More reconstructions results of i-MAE on ImageNet-1K *validation* images with different mixing coefficients α (listed on the left) from models pre-trained with the subordinate image I_s as the only target, 0.5 mask ratio, and with distillation.

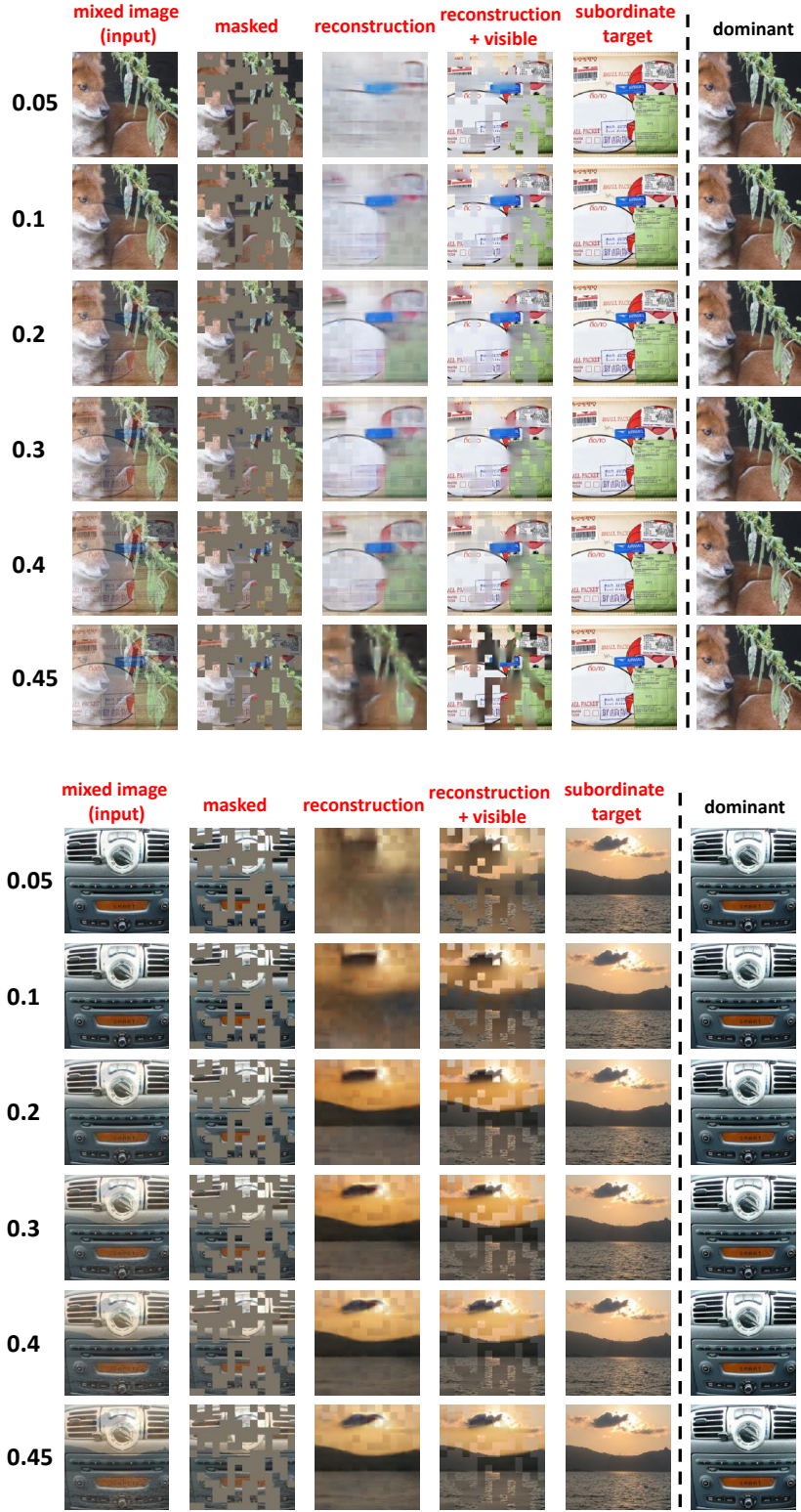


Figure 7: More reconstructions results of i-MAE on ImageNet-1K *validation* images with different mixing coefficients α (listed on the left) from models pre-trained with the subordinate image I_s as the only target, 0.5 mask ratio, and with distillation loss.

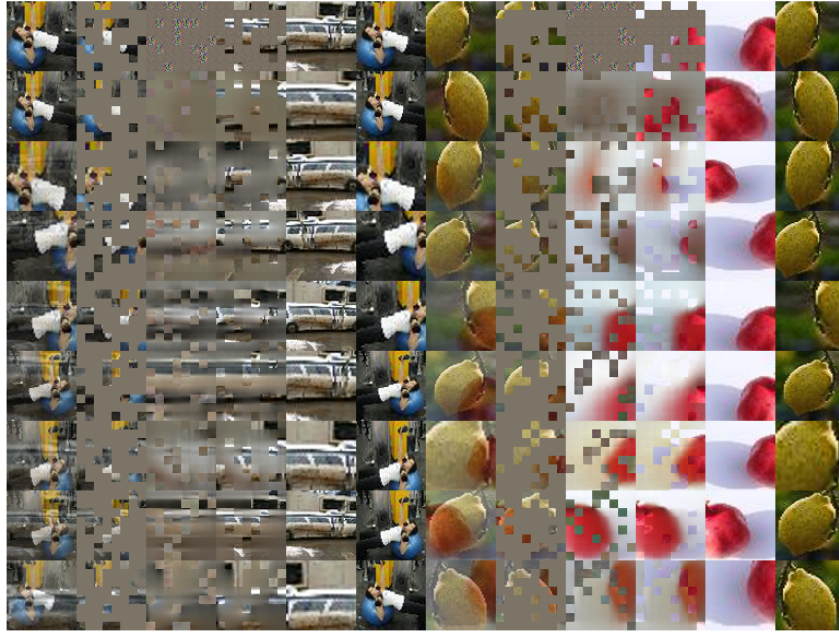


Figure 8: Uncurated Tiny-ImageNet reconstructions of different mix ratio, from 0.05 to 0.45, subordinate images.



Figure 9: Visual reconstructions of Tiny-ImageNet validation images using *semantics-controllable mixture* pre-trained i-MAE.

Table 5: Self-supervised pre-training configurations on CIFAR-10/100, Tiny-ImageNet, ImageNet-1K. For better visualizations, we use a 0.5 mask ratio on ImageNet.

Config	CIFAR-10/100	Tiny-ImageNet	ImageNet-1K
base learning rate	1.5e-4	1.5e-4	1e-3
batch size	4,096	4,096	4,096
Mask Ratio	0.75	0.75	0.5
optimizer	AdamW	AdamW	AdamW
optimizer momentum	0.9, 0.95	0.9, 0.95	0.9, 0.95
augmentation	None	RandomResizedCrop	RandomResizedCrop

Table 6: Pre-training Configurations of with the *semantics-controllable mixture* scheme.

Config	CIFAR-10/100	Tiny-ImageNet
Object mix range	0.0, 0.25, 0.5, 1.0	0.0, 0.25, 0.5, 1.0
Image mix ratio	$Beta(1.0, 1.0)$	$Beta(1.0, 1.0)$
base learning rate	1.5e-4	3.5e-4
batch size	4,096	4,096
Mask Ratio	0.75	0.75
optimizer	AdamW	AdamW
optimizer momentum	0.9, 0.95	0.9, 0.95
augmentation	None	RandomResizedCrop

Table 7: Finetune Classification Configurations.

Config	CIFAR-10/100	Tiny-ImageNet	ImageNet-1K
Object mix range	0.0 - 1.0	0.0 - 1.0	0.0, 0.25, 0.5, 1.0
Image mix ratio	$Beta(1.0, 1.0)$	$Beta(1.0, 1.0)$	0.8
base learning rate	1e-3	1e-3	1e-3
batch size	128	256	1,024
epochs	100	100	25
optimizer	AdamW	AdamW	AdamW
optimizer momentum	0.9, 0.999	0.9, 0.999	0.9, 0.999
augmentation	Mixup	Mixup, RandomResizedCrop	Mixup, RandomResizedCrop

Table 8: Linear Classification Configurations.

Config	CIFAR-10/100	Tiny-ImageNet	ImageNet-1K
Object mix range	0.0 - 1.0	0.0 - 1.0	0.0, 0.25, 0.5, 1.0
Image mix ratio	0.35	0.35	0.35
base learning rate	1e-2	1e-2	1e-2
batch size	128	256	1,024
epochs	200	200	25
optimizer	SGD	SGD	SGD
optimizer momentum	0.9, 0.999	0.9, 0.999	0.9, 0.999
augmentation	None	RandomResizedCrop	RandomResizedCrop

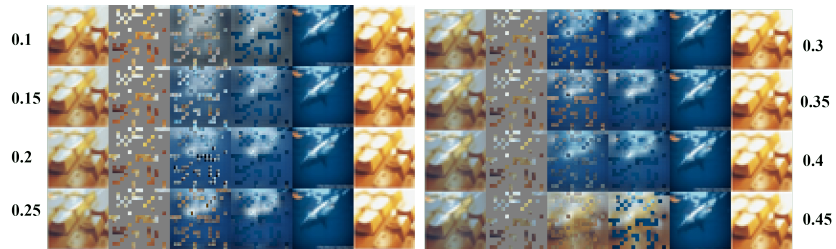


Figure 10: CIFAR-100 subordinate reconstruction of different ratios marked on the left and right side. Similarly, reconstructions at 0.45 are confused with the dominant image.



Figure 11: Uncurated reconstructions of CIFAR-100 validation images using *semantics-controllable mixture* from 0.0 (topmost) to 1.0 (bottom), in 0.1 intervals.