

FoMo-0D: A Foundation Model for Zero-shot Tabular Outlier Detection

Yuchen Shen
Carnegie Mellon University

yuchens3@andrew.cmu.edu

Haomin Wen
Carnegie Mellon University

haominwe@andrew.cmu.edu

Leman Akoglu
Carnegie Mellon University

lakoglu@andrew.cmu.edu

Reviewed on OpenReview: <https://openreview.net/forum?id=XCQzwpR9jE>

Abstract

Outlier detection (OD) has a vast literature as it finds numerous real-world applications. Being an unsupervised task, model selection is a key bottleneck for OD without label supervision. Despite a long list of available OD algorithms with tunable hyperparameters, the lack of systematic approaches for unsupervised algorithm and hyperparameter selection limits their effective use in practice. In this paper, we present **FoMo-0D**, a pre-trained **Foundation Model** for zero/0-shot OD on tabular data, which bypasses the hurdle of model selection altogether. Capitalizing on synthetic pre-training, **FoMo-0D** can directly predict the (outlier/inlier) label of test samples without parameter fine-tuning. Even more importantly, *it requires no labeled data, and no additional training or hyperparameter tuning when given a new task*. Extensive experiments on **57** real-world datasets against **26** baselines show that **FoMo-0D** is highly competitive; outperforming the majority of the baselines with no statistically significant difference from the 2nd best method. Further, **FoMo-0D** is efficient in inference time requiring only **7.7 ms** per sample on average, with at least **7x** speed-up compared to previous methods. To facilitate future research, our implementations for data synthesis and pre-training as well as model checkpoints are openly available at <https://github.com/A-Chicharito-S/FoMo-0D>.

1 Introduction

Outlier detection (OD) in tabular data finds numerous applications in critical domains such as security, environmental monitoring, finance, and medicine, to name a few. This popularity brings along a large literature with plethora of detection algorithms to choose from given a new OD task. These algorithms, however, exhibit several hyperparameters (HPs) that need careful tuning (Ma et al., 2023). Since most OD tasks are unsupervised¹, what makes effective detection notoriously difficult is unsupervised model selection (both algorithm and HP selection) in the absence of labels.

While deep learning has revolutionized many areas of machine learning (ML), it is not quite the case for OD. This is mainly because compared to classical methods, deep OD models (Pang et al., 2021) have many more HPs to which the detection performance is sensitive (Ding et al., 2022), making model selection even more challenging. While the recent success of large foundation models, pre-trained on massive amounts of data, offers new opportunities for zero-shot OD, thus far the most notable progress has been in NLP and computer vision (Brown et al., 2020; Touvron et al., 2023; Radford et al., 2021). This is thanks to the admirable

¹While supervised OD exists, unsupervised setting is preferred in most domains to detect novel, emergent anomalies.

Table 1: p -values of the one-sided Wilcoxon signed rank test, comparing FoMo-0D (with $D = 100$) to **top 10 baselines** with default hyperparameters (HPs), and **top 4^{avg}** baselines⁶ with **avg.** performance over varying HPs (denoted w/ ^{avg}) over All (57) datasets, those (42) w/ $d \leq 100$ and (46) w/ $d \leq 500$ dimensions, and those (47) excluding NLP, CV datasets. FoMo-0D shows **no statistically significant difference from the 2nd best model** (k NN, w/ $p = 0.106$) over All datasets, and none of the differences are significant when embedded, i.e., NLP and CV datasets are excluded, while it is comparable to ($p > \alpha$) or significantly better than ($p > 1 - \alpha$) all 26 original + 4^{avg} baselines over datasets w/ $d \leq 100$ (aligned w/ pretraining where $D = 100$) as well as on datasets w/ $d \leq 500$ (generalizing beyond pretraining). We use underline and **bold** to indicate $p < \alpha$ and $p > 1 - \alpha$. Rank, avg.'ed over all 57 datasets by AUROC. (setting: $D = 100$, $R = 500$, train/inference context size=5K, w/ quantile transform, $\alpha = 0.05$) (See Tables 17.1&17.2 for full results.)

	FoMo-0D	DTE-NP	kNN	ICL	DTE-C	LOF	CBLOF	Feat.Bag.	SLAD	DDPM	OCSVM	DTE-NP ^{avg}	kNN ^{avg}	ICL ^{avg}	DTE-C ^{avg}
$d \leq 100$	-	0.415	0.700	0.949	0.953	0.970	0.971	0.996	0.876	0.980	0.978	0.752	0.860	0.958	1.000
$d \leq 500$	-	0.220	0.569	0.827	0.894	0.960	0.968	0.994	0.910	0.960	0.979	0.607	0.756	0.846	1.000
All	-	<u>0.016</u>	0.106	0.462	0.454	0.585	0.750	0.823	0.759	0.901	0.895	0.112	0.315	0.670	1.000
All \ {NLP, CV}	-	0.164	0.476	0.757	0.832	0.934	0.945	0.988	0.867	0.938	0.965	0.515	0.683	0.777	1.000
Rank(avg)	11.886	7.553	9.018	10.851	11.36	12.316	13.342	13.386	12.982	14.061	13.851	9.079	11.105	12.991	22.263

quantity and quality of public text and image datasets. In comparison, public tabular OD benchmarks remain minuscule (Han et al., 2022; Zhao et al., 2021; Steinbuss & Böhm, 2021).

Recently, Prior-data Fitted Networks (PFNs) has marked a milestone in ML as a new approach to learning on tabular data (Müller et al., 2022). The core idea is to compute a posterior predictive distribution (PPD) for a test point given training data as context. To approximate the PPD, a Transformer (Vaswani et al., 2017) is pre-trained on a large set of synthetic datasets drawn from pre-defined data priors. At inference, the pre-trained PFN is fed with test samples along with some training samples as context for zero-shot prediction, requiring no parameter fine-tuning or model selection on new datasets. Variants of PFN are shown to match tree-based models in performance on small classification datasets (Hollmann et al., 2023) as well as time series forecasting (Dooley et al., 2023).

In this paper, we capitalize on these ideas and introduce FoMo-0D: a prior-data fitted **F**oundation **M**odel for zero/0-shot OD. Once pre-trained on synthetic datasets, FoMo-0D unlocks zero-shot OD on a new dataset where the (unlabeled) input data is fed only as context. As such, FoMo-0D bypasses not only model (parameter) training, but more importantly, the nontrivial task of unsupervised model (algorithm and HP) selection without labeled data. Figure 1 illustrates the new FoMo-0D paradigm versus the typical OD setting. To our knowledge, FoMo-0D is the first pre-trained foundation model for tabular OD.

In designing FoMo-0D, we use Gaussian mixture models as a simple yet effective tabular data prior for inlier data distributions (Hollmann et al., 2023; Zhao et al., 2021), which are also employed to simulate outlier types common in the real world; namely, local and global subspace outliers (Steinbuss & Böhm, 2021). While the data prior can be extended to comprise more complex data distributions (Hollmann et al., 2023) (e.g. Bayesian Neural Networks (Neal, 2012) and Structural Causal Models (Pearl, 2009)), and additional outlier types can be included (e.g. dependency, contextual, etc.), as we show in the experiments, even with its relatively straightforward prior, FoMo-0D achieves remarkable performance: As shown in Table 1 FoMo-0D, which is pre-trained on datasets with $d \leq 100$ dimensions, shows no statistically significant difference from all 26 state-of-the-art baselines (all p -values > 0.4) on 42 benchmark datasets with dimensionality $d \leq 100$ (aligned with pre-training), while our method consistently ranks among the top and outperforms a majority of the baselines with p -value > 0.95 . (See Appendix Tables 17.1&17.2 for full results.) The results remain consistent on 46 benchmarks with $d \leq 500$ dimensions. FoMo-0D is also competitive across all 57 datasets, effectively generalizing beyond its pre-training distributions, with no statistically significant difference from the 2nd best baseline. When intrinsically non-tabular datasets are removed (i.e., datasets from NLP, CV domains embedded with pre-trained encoders), there is no statistical evidence to suggest a performance difference between FoMo-0D and any baseline on the remaining 47 datasets. Further, FoMo-0D takes a mere 7.7 ms to infer a test sample on average with no extra training or tuning overhead on the new dataset. We summarize the main contributions of our work as follows.

- **A Foundation Model for Tabular OD:** We present FoMo-0D, *the first foundation model for zero-shot OD* on unseen tabular datasets, with no additional training or hyperparameter tuning, backed by Transformer-based in-context learning (ICL), synthetic data pre-training, and feed-forward inference.
- **Model Selection Made Obsolete:** FoMo-0D is designed for zero-shot inference given a new dataset, fully abolishing not only model training on the new dataset, but also the notorious task of algorithm selection and hyperparameter tuning in the absence of labeled data.
- **Scalable Pre-training:** To enable pre-training on many large datasets, we propose (i) a new mechanism to reduce sample-to-sample attention from quadratic to linear time—enabling larger datasets, as well as (ii) on-the-fly data synthesis through data transformations—enabling more diverse datasets in less time.
- **Fast OD at Inference:** Given a new dataset, FoMo-0D bypasses both model training and selection, both of which can be slow for modern deep OD models with many hyper/parameters. Rather, it takes fraction of a second to label a test point through a single forward pass. Such speedy inference also unlocks the potential for deploying FoMo-0D in real time on data streams.
- **Effectiveness:** On a large benchmark of **57** datasets (Han et al., 2022) from diverse domains and against **26** baselines ranging from classical to modern OD models (Livernoche et al., 2024), FoMo-0D outperforms the majority of the baselines, with no statistical evidence for performance difference from the *2nd* best baseline, while operating fully zero-shot on real-world datasets out-of-the-box.
- **New directions:** As the first foundation model for OD, FoMo-0D presents a paradigm shift in how to perform OD in practice, while offering new directions to explore as well as open questions to investigate. From the algorithmic perspective, how can we understand what (OD) algorithm, if any, has the Transformer- and ICL-based FoMo-0D learned? How can we interpret (mechanistically or otherwise) how the pretraining achieves zero-shot and out-of-distribution (OOD) generalization? From a data perspective, what prior distributions for pretraining are suitable for downstream real-world tasks? What is the role of prior data diversity and complexity in the generalization ability of the model? In summary, FoMo-0D paves the path towards a better understanding of ICL as well as building more powerful future foundation models for OD.

2 Problem and Preliminaries

2.1 Outlier Detection Problem and Setting

Outlier detection (OD) methods can be categorized based on the availability of labeled data. In supervised OD, the task is similar to binary classification with imbalanced classes (as outliers typically make up only a small portion of the overall data). The more difficult unsupervised setting assumes that the “contaminated” training data contains both inliers and outliers, but without any labels. This is also the transductive setting where training and test data are the same. The one-class setting lies between these two extremes, where the “clean” training data consists only of inliers, while the unlabeled test data contains the outliers to be detected. Here, training and test sets are disjoint and thus the setting is inductive. Note that there exist **no labeled outliers** in both settings, making model selection challenging.

We remark that some OD literature refers to the latter setting as semi-supervised OD, which is a misnomer from the supervised ML perspective where semi-supervised classification assumes the presence of some labeled instances from **all** classes in the training data. In the rest of text, we adopt the terminology unsupervised OD for both settings, and specify “clean” inlier-only versus “contaminated” mixed training data to differentiate them. Then, our work considers the unsupervised OD problem under “clean” inlier-only training data.

Formally, let $\mathcal{D}_{\text{in}} = \{(\mathbf{x}_1, y_1) \dots, (\mathbf{x}_n, y_n)\}$ denote the inlier-only input data $\mathbf{x}_i \in \mathbb{R}^d$, where $y_i = 0 \ \forall i \in [n]$, and $\mathcal{D}_{\text{test}}$ depicts the unlabeled test data comprising both inliers and outliers. Note that $\mathcal{D}_{\text{in}} \cap \mathcal{D}_{\text{test}} = \emptyset$, i.e. train/test split is disjoint. The task is to assign labels to $\mathbf{x}_i \in \mathcal{D}_{\text{test}}$ given the inlier-only input $\mathcal{D}_{\text{in}}$.

2.2 Background on Prior-data Fitted Networks

Posterior Predictive Distribution (PPD): In the Bayesian framework for supervised learning, the prior defines a hypothesis space Φ which expresses our beliefs about the data distribution before seeing any data.

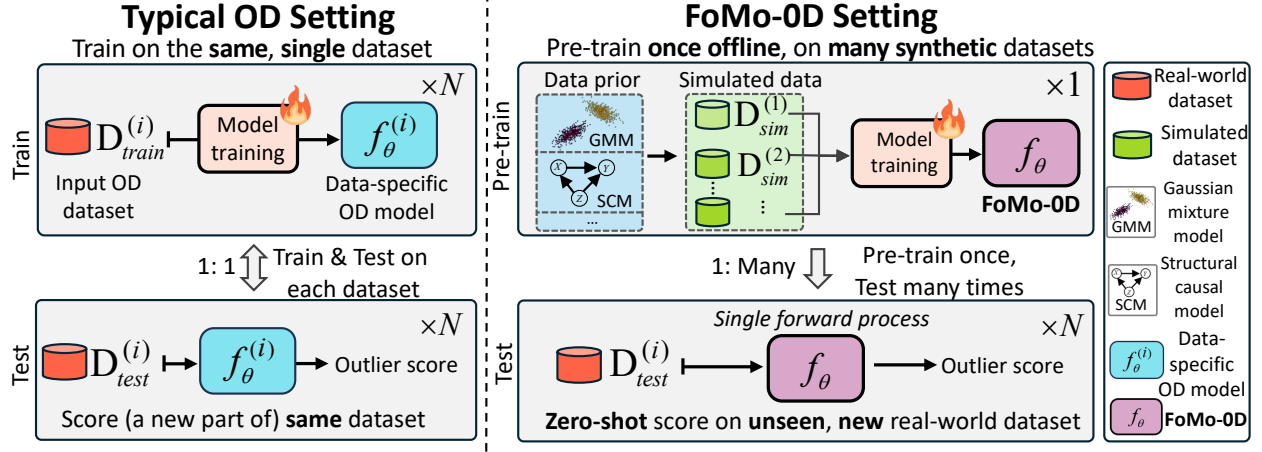


Figure 1: (best in color) Comparison of typical OD vs. the FoMo-OD settings. Given a new unlabeled OD dataset, FoMo-OD not only eliminates the need for model (parameter) training, but most importantly, also abolishes the onerous task of unsupervised model selection (algorithm and hyperparameters).

Each hypothesis $\phi \in \Phi$ describes a mechanism by which the data is generated. The posterior predictive distribution $p(\cdot | \mathbf{x}_{\text{test}}, \mathcal{D}_{\text{train}})$ provides a framework for making prediction on new, unseen test data $\mathbf{x}_{\text{test}}$, conditioned on observed training data $\mathcal{D}_{\text{train}} = \{(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_n, y_n)\}$. Based on Bayes' Theorem, the PPD can be derived by the integration over the space of hypotheses Φ :

$$p(y_{\text{test}} | \mathbf{x}_{\text{test}}, \mathcal{D}_{\text{train}}) = \int_{\Phi} p(y_{\text{test}} | \mathbf{x}_{\text{test}}, \phi) p(\mathcal{D}_{\text{train}} | \phi) p(\phi) d\phi, \quad (1)$$

where $p(\phi)$ denotes the prior probability and $p(\mathcal{D} | \phi)$ is the likelihood of the data $\mathcal{D}$ given ϕ .

PFNs and PPD Approximation: As obtaining the above PPD is generally intractable, Prior-data Fitted Networks (PFNs) are proposed to approximate the PPD (Müller et al., 2022). Unlike traditional machine learning models that are trained directly on observed datasets, PFNs are pre-trained on simulated datasets that are generated according to a prior distribution. Specifically, it contains the pre-training and inference stages described as the following.

Pre-training on synthetic data. Massive synthetic datasets are generated for the pre-training stage, by first sampling a hypothesis (i.e., the generating mechanism) $\phi \sim p(\phi)$, and then sampling a dataset $\mathcal{D} \sim p(\mathcal{D} | \phi)$. For training, each dataset $\mathcal{D}$ can be split as $\mathcal{D}_{\text{test}} \subset \mathcal{D}$ and $\mathcal{D}_{\text{train}} = \mathcal{D} \setminus \mathcal{D}_{\text{test}}$. Thus, the PFN with parameters θ can be optimized by making predictions on data points in $\mathcal{D}_{\text{test}}$. For a test point $(\mathbf{x}_{\text{test}}, y_{\text{test}}) \in \mathcal{D}_{\text{test}}$, the training loss is formulated as follows.

$$\mathcal{L} = \mathbb{E}_{(\{\mathbf{x}_{\text{test}}, y_{\text{test}}\}) \cup \mathcal{D}_{\text{train}} \sim p(\mathcal{D})} [-\log q_{\theta}(y_{\text{test}} | \mathbf{x}_{\text{test}}, \mathcal{D}_{\text{train}})] . \quad (2)$$

The above loss can also be interpreted as minimizing the expected KL divergence between $p(\cdot | \mathbf{x}, \mathcal{D})$ and $q_{\theta}(\cdot | \mathbf{x}, \mathcal{D})$ (Müller et al., 2022). In practice, a PFN model q_{θ} is typically implemented by a Transformer-based architecture (Vaswani et al., 2017), which takes $(\mathbf{x}_{\text{test}}, \mathcal{D}_{\text{train}})$ as input, where $\mathbf{x}_{\text{test}} \in \mathcal{D}_{\text{test}}$ and $\mathcal{D}_{\text{train}}$ contains an arbitrary number of instances. The output is the conditional class probabilities for $\mathbf{x}_{\text{test}}$. As the whole training set $\mathcal{D}_{\text{train}}$ is passed as input/context to the Transformer, it learns to predict class labels through sample-to-sample attention.

Inference on real-world data. In the inference stage, a fresh real-world dataset $\mathcal{D}_{\text{train}}$ and some test instance $\mathbf{x}_{\text{test}}$ are fed into the (frozen) pre-trained model, which computes the PPD $q_{\theta}(\cdot | \mathbf{x}_{\text{test}}, \mathcal{D}_{\text{train}})$ in a single forward pass. Importantly, PFNs do not require gradient-based parameter tuning on new datasets, where prediction is delivered *in less than a second* (Hollmann et al., 2023).

In summary, PFNs are trained once, and can be used many times for zero-shot inference on new datasets with different characteristics. The main benefit is that **no training or tuning** is required at the inference stage. This type of learning ability is also termed as in-context learning (ICL) (Xie et al., 2021), which was

shown to be effective for various tasks in languages (Brown et al., 2020). In fact, ICL with PFNs is recently shown to be a promising paradigm for supervised classification on tabular datasets (Hollmann et al., 2023).

3 FoMo-OD: A Foundation Model for Zero-shot Tabular Outlier Detection

Inspired by PFNs (Müller et al., 2022; Hollmann et al., 2023; Dooley et al., 2023), we propose FoMo-OD for zero-shot OD, which is pre-trained on large-scale synthetic OD datasets for zero-shot detection at inference time. FoMo-OD eliminates the need for model training on a new dataset and for model selection (both algorithm and HPs), which is difficult without any labeled data. The new FoMo-OD paradigm (right) versus the typical OD setting (left) is illustrated in Figure 1.

In the following we describe our OD data prior, training on prior-simulated datasets, inference on new datasets, and model architecture and improvements for scalability.

3.1 Designing a Data Prior for Outlier Detection

Foundation models benefit from massive amounts of datasets available for pre-training, along with high-capacity model architectures, however, the quantity (and quality) of publicly available tabular OD datasets is minuscule, compared to the massive size of open-domain text corpora. Even with large quantities of data, Ansari et al. (2024) show that using synthetic data in combination with real-world data improves the overall zero-shot performance for time-series foundation models. Hence, we design a new data prior from which we simulate numerous OD datasets for pre-training FoMo-OD.

Ideally, the data prior should reflect distributions as general and diverse as seen in the real world, however, “finding a prior supporting a large enough subset of possible [data generating] functions isn’t trivial” (Nagler, 2023). Surprisingly, our results show that a straightforward, simple-to-implement data prior is sufficient to achieve remarkable performance.

Inlier synthesis: We simulate inliers by drawing from a Gaussian Mixture Model (GMM) with m -clusters in d -dimensions, with centers $\mu_{jk} \in [-5, 5]$, $j \in [m]$, $k \in [d]$ and *diagonal*² Σ_j with entries in $(0, 5]$. We create different GMMs with varying $m \leq M$ and $d \leq D$ chosen uniformly at random from $[M]$ and $[D]$, respectively. From each GMM, we draw a set of S inliers, defined as instances within the 90th percentile of the GMM.

Outlier synthesis: Following Han et al. (2022), we generate subspace outliers by first drawing a subset of dimensions $\mathcal{K}$ at random, for $|\mathcal{K}| \leq d$, and then generate S points from the “inflated” GMM, which shares the same centers μ_j ’s with the original GMM but with the inflated (diagonal) covariances $5 \times \Sigma_{j, kk}$ ’s for $k \in \mathcal{K}$. Outliers are defined as points sampled outside the 90th percentile of the original GMM, which are labeled based on their Mahalanobis distances (see Property B.6 in the Appendix).

Specifically, we simulate datasets containing $2S = 10,000$ samples (half inlier, half outlier) from the two corresponding GMMs (original and inflated) with up to $M = 5$ clusters and up to $D = 100$ dimensions. Example 2- d synthetic datasets are illustrated in Appendix A.

Remarks: Our model is not trained on **any** real-world data but rather, on purely synthetic data (although future work can combine existing benchmark OD datasets with synthesized data, as was done by Ansari et al. (2024) for time series). While we have intended to extend our preliminary attempt toward designing a sophisticated data prior for OD, we found (to our surprise) that even with a basic, GMM-based prior, FoMo-OD generalizes remarkably well to real-world OD datasets downstream³, outperforming numerous SOTA baselines. Therefore, we present FoMo-OD with this simple prior to showcase the prowess of PFNs for OD. We leave as future work the exploration of other priors (Hollmann et al., 2023) and other outlier types (contextual, dependency, etc. (Steinbuss & Böhm, 2021)), the impact of different priors on performance, as well as prior mixture composition to further improve performance.

²In early experiments, we found no difference in test performance on synthetic datasets between using diagonal vs. non-diagonal Σ , yet, it is easier to invert diagonal Σ for data synthesis.

³We refer to Appendix G.2 and Figure 16 for an exploration of FoMo-OD performance and GMM goodness of fit on real-world OD datasets.

3.2 (Pre)Training and Inference

Model (Pre)Training (Once, Offline): FoMo-0D is a Prior-data Fitted Network (PFN, see Section 2.2) based on the Transformer architecture. In the synthetic prior-data fitting phase, it is trained on datasets drawn from our OD data prior for tabular data introduced in Section 3.1. Each dataset is simulated from a different GMM configuration based on randomly drawn parameters, and consists of varying number of training samples and dimensions to capture the diversity in real-world tabular datasets. Details are outlined in Algorithm 1 in Appendix C, and described as follows.

At each time, we first draw a hypothesis (i.e. GMM configuration) uniformly at random, that is, $\phi = \{d \in [D], m \in [M], \{\mu_j\}_{j=1}^m \in [-5, 5]^d, \{\Sigma_j\}_{j=1}^m; \text{diag}(\Sigma_j) \in [-5, 5]^d\}$, and then generate a synthetic dataset $\mathcal{D} = \{\mathcal{D}_{\text{in}}, \mathcal{D}_{\text{out}}\}$ containing synthetic inlier and outlier samples from the drawn hypothesis and its variance-inflated variant, respectively.

We optimize FoMo-0D’s parameters θ to make predictions on $\mathcal{D}_{\text{test}} = \{\mathcal{D}_{\text{test}}^{\text{in}}, \mathcal{D}_{\text{test}}^{\text{out}}\}$, conditioned on the inlier-only training data $\mathcal{D}_{\text{train}} \subset \mathcal{D}_{\text{in}}$ based on the cross-entropy loss (see Eq. (2)). During training, $\mathcal{D}_{\text{test}}$ contains a *balanced* number of inlier and outlier samples, where $\mathcal{D}_{\text{test}}^{\text{in}} = \mathcal{D}_{\text{in}} \setminus \mathcal{D}_{\text{train}}$, and $\mathcal{D}_{\text{test}}^{\text{out}} \subset \mathcal{D}_{\text{out}}$ contains an equal number of samples as $\mathcal{D}_{\text{test}}^{\text{in}}$. To vary the training data size, we subsample $\mathcal{D}_{\text{train}}$ of randomly drawn size $n \in [n_L, n_U]$, where n_L and n_U denote the lower and upper bounds. In our implementation, we use $n_L = 500$, and $n_U = 5,000$.

FoMo-0D is trained on 200,000 batches (200 epochs $\times$ 1,000 steps/epoch) of $B = 8$ generated datasets in each batch. While this pre-training phase can be expensive, it is done *only once, offline*. Moreover, we introduce several scalability improvements to speed up pre-training, as discussed later in Section 3.3. Full details on the training and implementation of FoMo-0D are given in Appendix C.

Zero-shot Inference (on Unseen/New Dataset): At inference, the pre-trained FoMo-0D can be employed on any unseen real-world dataset. Specifically, for a new unsupervised OD task with inlier-only training data $\mathcal{D}_{\text{train}}$ and mixed test data $\mathcal{D}_{\text{test}}$, feeding $\langle \mathcal{D}_{\text{train}}, \mathbf{x}_{\text{test}} \rangle$ as input to FoMo-0D (for each $\mathbf{x}_{\text{test}} \in \mathcal{D}_{\text{test}}$ separately) yields the PPD $q_\theta(y|\mathbf{x}_{\text{test}}, \mathcal{D}_{\text{train}})$ in a *single forward pass*. As such, FoMo-0D performs model “training” and prediction simultaneously at test time. In fact, as the training data is passed as context, FoMo-0D leverages in-context learning (ICL) (Xie et al., 2021; Garg et al., 2022) for inference.

Remarks: The **key** contribution of FoMo-0D goes beyond eliminating the need for model training for a new dataset, it *renders model selection an obsolete concern for OD*. In other words, a practitioner with a new detection task no longer needs to choose an OD model to train or grapple with tuning any hyperparameters of the said model. Further, the speedy, easily parallelizable inference (for *less-than-a-second* per test sample) is the “icing on the cake”. Figure 1 (right) illustrates (top) pre-train and (bottom) test phases of FoMo-0D, where the pre-trained FoMo-0D is reused during inference on new datasets directly, unlocking zero-shot OD.

3.3 Architecture and Scalability

Architecture and sample-to-sample attention: Like existing PFNs, FoMo-0D is based on the Transformer (Vaswani et al., 2017), encoding each sample’s feature vector as a fixed size token through a linear embedding layer (see Appendix C.2), and allowing token representations to attend to each other, hence enabling sample-to-sample attention. We also adopt the three customizations from TabPFN (Hollmann et al., 2023), which (1) computes self-attention among all the training samples and only cross-attention from test samples to the training samples, (2) enables varying feature dimensionality by zero-padding, and (3) randomly permutes input samples while omitting positional encodings to achieve model invariance in the dataset.

Given $\mathcal{D}_{\text{train}} = \{\mathbf{x}_1, \dots, \mathbf{x}_n\}$, each self-attention layer outputs n embeddings $\{\mathbf{z}_i\}_{i=1}^n$; where the i -th token is mapped via linear transformations to a key $\mathbf{k}_i$, query $\mathbf{q}_i$ and value $\mathbf{v}_i$, where the i -th output is computed as

$$\mathbf{z}_i = \sum_{j=1}^n \text{softmax}(\{\langle \mathbf{q}_i, \mathbf{k}_j \rangle\}_{j=1}^n)_j \cdot \mathbf{v}_j. \quad (3)$$

The sample-to-sample attention is intriguing from the perspective of OD: many classical OD algorithms (Aggarwal, 2013) are based on nonparametrics; in particular, they leverage the distances to the k *nearest* neighbors (k NNs) of a point to compute its outlierness, where k is a critical hyperparameter. One can think

of FoMo-0D as mimicking non-parametric models but by using parametric attention mechanisms. Interestingly, PFNs are much more robust and flexible than k NN based OD approaches, for (1) sample-to-sample relations are not pre-specified but rather learned through attention weights, and thus (2) they are not limited to just the nearest neighbors but rather can *learn which* training points are worth attending to, and (3) as attention is dataset-wide across all points, there is no need for specifying a cut-off HP value like k , to which most k NN based OD techniques are sensitive to (Aggarwal & Sathe, 2015; Campos et al., 2016; Goldstein & Uchida, 2016; Ding et al., 2022). We present analyses on sample-to-sample attention in Appendix E.

To seize the power of scale, we incorporate a scalable architecture and data synthesis into our design to benefit pre-training and inference, as we describe next. The scale-up unlocks a larger context size for FoMo-0D, enabling pre-training and inference on larger datasets with fast speed.

Scaling up attention with “routers”: The $\mathcal{O}(n^2)$ quadratic sample complexity at pre-training presents an obstacle for achieving high performance at inference, as it limits pre-training to relatively small training datasets, and degenerates in-context learning that typically benefits from longer context (Xie et al., 2021).

Toward a high-performance model, we scale up FoMo-0D’s attention via the “router mechanism” of Zhang & Yan (2023). As shown in Figure 2, the main idea is to learn a small number ($R \ll n$) of “routers”, which gather information from all n samples and then distribute the information back to the n output embeddings, in effect, reducing complexity from $\mathcal{O}(n^2)$ to $\mathcal{O}(2Rn) = \mathcal{O}(n)$. This design allows FoMo-0D to **scale linearly** with respect to both dimensionality d and dataset size n in pre-training as well as during inference.

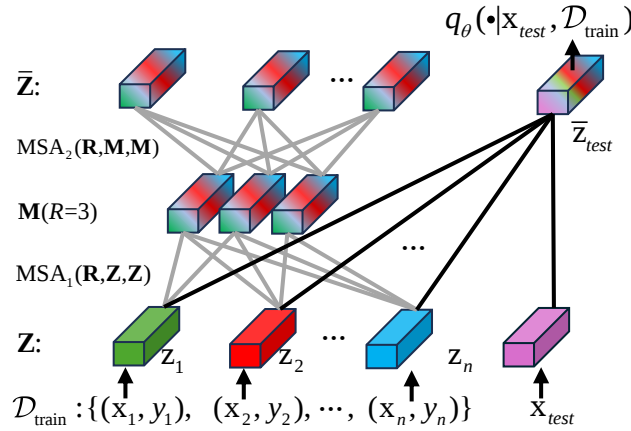


Figure 2: FoMo-0D architecture employs the “router mechanism” for scalable attention.

Concretely, the representatives first aggregate information from all samples by serving as the query in the multi-head self-attention (MSA):

$$\mathcal{M} = \text{MSA}_1(\mathbf{R}, \mathbf{Z}, \mathbf{Z}) , \quad (4)$$

where $\mathbf{R} \in \mathbb{R}^{R \times d}$ depicts the *learnable* vector array of representatives and $\mathcal{M}$ denotes the aggregated messages. Then, the routers distribute the received information among samples by using the sample embeddings as query and the aggregated messages as both key and value:

$$\hat{\mathbf{Z}} = \text{MSA}_2(\mathbf{Z}, \mathcal{M}, \mathcal{M}) . \quad (5)$$

Finally, we obtain $\bar{\mathbf{Z}} = \text{LayerNorm}(\hat{\mathbf{Z}} + \mathbf{Z})$ after layer normalization. Note that the test samples only attend to the training samples’ embeddings, computed in the described manner across layers, and are finally fed into the prediction head to estimate the PPD at the output layer.

Scaling up (pre)training data synthesis with linear transforms: Besides the scalability challenge associated with architecture/attention, another computational challenge in pre-training FoMo-0D arises from drawing samples from the data prior, which requires considerable time, especially in high dimensions⁴,

⁴This is because the inverse of the $(d \times d)$ covariance matrix plays a crucial role in the process of drawing samples from GMMs, which has $\mathcal{O}(d^3)$ time complexity. (It is also the reason why diagonal Σ_j ’s are favored in our data prior.) In addition, Mahalanobis distance for labeling inliers/outliers also requires the inverse.

provided the large number of datasets we sample (specifically, we utilize a batch size of 8 datasets over 1,000 steps each for 200 epochs).

To give an idea, sampling a dataset with $n = 10,000$ points in $d = 100$ dimensions using 10 CPUs in parallel takes ≈ 0.4 seconds (see Appendix Figure 7). Across 200 training epochs with 1,000 steps each, it adds up to more than 177 hours just to generate 1.6 million datasets on-the-fly. Of course, one can trade storage with compute-time by generating all these datasets apriori via massive parallelism. Nevertheless, synthetic data generation demands considerable time (and/or storage).

To scale up data synthesis, FoMo-0D employs two distinct strategies. **First**, we propose *reuse at epoch level*: that is, one can reuse the same 8K (8×1000) unique datasets at every epoch, or in general, the same $8K \times P$ datasets periodically at every P epochs. A larger P would lead to more diversity in terms of the overall pre-training data used.

Second, we propose *reuse at dataset level via transformation*: that is, having generated one unique dataset $\mathbf{X} \in \mathbb{R}^{n \times d}$ from a GMM, we propose a linear transform $T(\mathbf{x})$ of the form $\mathbf{W}\mathbf{x} + \mathbf{b}$ for randomly drawn parameters $\mathbf{W} \in \mathbb{R}^{d \times d}$ and $\mathbf{b} \in \mathbb{R}^d$ (see Appendix B.1).⁵ This simple yet efficient transformation creates a new dataset, akin to one being drawn from another GMM with centers $T(\boldsymbol{\mu}_j) = \mathbf{W}\boldsymbol{\mu}_j + \mathbf{b}$ and covariance $T(\boldsymbol{\Sigma}_j) = \mathbf{W}\boldsymbol{\Sigma}_j\mathbf{W}^T$, $\forall j \in [m]$. Note that we do not actually materialize these parameters but only transform the dataset. As we show in the following, such transformations preserve the Mahalanobis distances as well as the percentile thresholds for labeling points as inlier/outlier. Details and proofs are given in Appendix B.

Lemma 3.1. *Linear transform T with invertible $\mathbf{W}$ on $\mathcal{G}_m^d$ preserves Mahalanobis distances.*

Lemma 3.2. *Linear transform T with invertible $\mathbf{W}$ on $\mathcal{G}_m^d$ preserves the percentiles of the GMM.*

The implication of these lemmas is that a linear transformation of a dataset from a GMM retains the identity of the inliers and outliers, i.e. no relabeling is required. Moreover, notice that as a byproduct we obtain a transformed dataset as though it is drawn from a GMM with a *non-diagonal* covariance matrix which, besides the time savings, offers a slightly more complex data prior.

To reach 8K unique datasets for each epoch, we first generate 500 datasets from different GMMs (with varying configurations), then employ 15 different linear transformations to each dataset by varying $\mathbf{W}$ and $\mathbf{b}$. Drawing each $(\mathbf{W}, \mathbf{b})$ takes ≈ 0.02 seconds, while the matrix-matrix product of $\mathbf{X}$ ($n \times d$) and $\mathbf{W}$ ($d \times d$) takes negligible time (for $d \leq 100$). Thus, obtaining a transformed dataset offers $20\times$ speed-up compared to generating one (0.02 vs. 0.4 seconds).

4 Experiments

4.1 Setup

We present the experiment setup briefly, including data synthesis, real-world datasets, baselines, metrics and HPs. For more details, we refer to Appendix D.

Pre-training Dataset Synthesis: During pre-training, we generate unique GMM datasets by first drawing a configuration, including dimensionality $d \in [D]$, number of components $m \in [M]$, centers $\{\boldsymbol{\mu}_j\}_{j=1}^m$ (each $\boldsymbol{\mu}_j \in [-5, 5]^d$) and covariances $\{\boldsymbol{\Sigma}_j\}_{j=1}^m$ ($\text{diag}(\boldsymbol{\Sigma}_j) \in [-5, 5]^d$). We set $M = 5$ and vary $D \in \{20, 100\}$ to study pre-training with relatively small and high dimensional datasets, respectively. We synthesize inliers and outliers described in Section 3.1.

Real-world Benchmark Datasets: While pre-training is purely on synthetic datasets, we evaluate FoMo-0D on **57** real-world datasets from ADBench (Han et al., 2022) (see Table 20 in Appendix J). Following Livernoche et al. (2024), we use 5 train/test splits of each dataset via different seeds and report mean performance and standard deviation. Note that the baselines require model re-training and inference for each $\mathcal{D}_{\text{train}}/\mathcal{D}_{\text{test}}$ split, while FoMo-0D uses the splits only for inference as $\mathcal{D}_{\text{train}}$ is passed as context.

⁵In practice, we apply the linear transform on the subspace of inflated features only, wherein inliers and outliers are defined, which remains to be a multi-variate GMM.

Baselines: We compare FoMo-0D against **26** baselines, from classical/shallow methods to modern/deep models. The baselines are imported from one of the latest papers that proposed the SOTA diffusion-based OD model, DTE (Livernoche et al., 2024), and three variants; DTE-C, DTE-IG, DTE-NP. As such, the long list of baselines we compare to constitutes one of the most comprehensive in the literature. We refer to the original paper for more details.

Model Implementation: We train our final model for 200,000 steps with a batch size of 8 datasets. That is, FoMo-0D is trained on 1,600,000 synthetically generated datasets. This training takes about 25 hours on 1 GPU (Nvidia RTX A6000). Each dataset had a fixed size of 10,000 samples, with $|\mathcal{D}_{\text{train}}| \in [n_L = 500, n_U = 5000]$, and the rest as $\mathcal{D}_{\text{test}}$ with balanced number of inliers and outliers. Other details of FoMo-0D, including the training algorithm, model architecture, data synthesis and reuse, and hardware are in Appendix C.

Metrics and Hypothesis Testing: Detection performance is w.r.t. 3 widely-used metrics for OD: AUROC; area under ROC curve, AUPR; area under Precision-Recall curve, and F1 score; using threshold at the true number of outliers in the test data (varies by dataset) Livernoche et al. (2024).

To compare different methods on ADBench, we compute their **rank** on each dataset (lower is better), and present **average rank** across datasets. This is an alternative to the average metric (e.g. AUROC), which is not meaningful when tasks vary widely in terms of their difficulties.

In addition, we perform significance tests to compare two methods statistically, using the one-sided paired Wilcoxon signed rank test (Demšar, 2006) between FoMo-0D and a baseline based on the performances across all datasets, with the alternative hypothesis suggesting the “baseline-minus-FoMo-0D” performance gap is greater than zero. We consider results to be significant at $\alpha = 0.05$ following convention.

Hyperparameters (HPs): Importantly, Livernoche et al. (2024) picked for each baseline the best-performing set of HPs as recommended by the authors in their original paper. As for their own DTE, which behaves similarly to k NN, they use $k = 5$ and set the *same* k for the k NN baseline (Ramaswamy et al., 2000) to be consistent. However, it is well known that k NN is sensitive to the value of k (Aggarwal & Sathe, 2015), and so are many other OD models to their respective HPs (Campos et al., 2016; Goldstein & Uchida, 2016; Zhao et al., 2021; Ding et al., 2022).

Therefore, besides comparing FoMo-0D with the 26 baselines in Livernoche et al. (2024), respectively for AUROC, F1, and AUPR (Livernoche et al., 2024), we also compare to the **top-4**⁶ best-performing baselines (in order: DTE-NP, k NN, ICL, and DTE-C) on their *average* performance across a list of different HP settings. Such an approach reflects their *expected* performance under HP values selected at random, in the absence of any other prior knowledge, as recommended by Goldstein & Uchida (2016) “*to get a fair evaluation when comparing [OD] algorithms*”. We annotate the method name with ^{avg} for the version with performance averaged over varying HPs. The detailed list of HP values for each top baseline is given in Appendix D.4. Overall, we compare FoMo-0D to 30 baselines; 26 from Livernoche et al. (2024) and ^{avg} of the top-4.

4.2 Results

Detection performance: Table 1 presented the comparison of FoMo-0D w/ $D = 100$ to all baselines w.r.t. average rank across datasets as well as pairwise Wilcoxon signed rank tests based on AUROC, and **we present full results on all datasets and all metrics in Appendix I**. We find that among 30 baselines and 2 variants of FoMo-0D (w/ $D = 100$ and $D = 20$), FoMo-0D w/ $D = 100$ performs as well as the *2nd* best model (k NN with default HP; $k = 5$) on all datasets. While DTE-NP outperforms FoMo-0D with author-recommended $k = 5$, we find that DTE-NP^{avg} is on par with FoMo-0D.

In our tests, $p > \alpha = 0.05$ implies no statistical evidence for performance difference between two methods. FoMo-0D w/ $D = 100$ performs statistically no different from **all** baselines on datasets with $d \leq 100$ (i.e., “at its own game” when pre-training data dimensions align with real-world datasets), while it outperforms the majority of baselines (where $p > 1 - \alpha$). These results continue to hold on datasets with $d \leq 500$.

⁶ To rank the 26 baselines, we compute the 26×26 p -values of the pairwise Wilcoxon signed rank test (see Appendix Figure 24), and order them by their mean p -value against other baselines.

Table 2: p -values of the one-sided Wilcoxon signed rank test, comparing FoMo-0D (w/ $D = 20$) to **top 10** baselines with default HPs, and **top 4^{avg}** baselines⁶ with **avg.** performance over varying HPs (denoted w/ ^{avg}) over All (57) datasets, those (24) w/ $d \leq 20$ and (38) datasets w/ $d \leq 50$ dimensions. Although pretrained on datasets w/ small $D = 20$, FoMo-0D shows **no statistically significant difference from the top 3rd baseline** (ICL, w/ $p = 0.089$) over All datasets, while it outperforms (w/ $p > 1 - \alpha$) the top 5th (LOF) and onward baselines over datasets w/ $d \leq 20$ (aligned w/ pretraining where $D = 20$) and on datasets w/ $d \leq 50$ (generalizing beyond pretraining). We use underline and **bold** to indicate $p < \alpha$ and $p > 1 - \alpha$. Rank is avg.'ed over all 57 datasets, where methods are ranked on each dataset w.r.t. AUROC. (experiment setting: $D = 20$, $P = 50$, $R = 500$, train/inference context size=5K, no data transformation, $\alpha = 0.05$)

	FoMo-0D	DTE-NP	kNN	ICL	DTE-C	LOF	CBLOF	Feat.Bag.	SLAD	DDPM	OCSVM	DTE-NP ^{avg}	kNN ^{avg}	ICL ^{avg}	DTE-C ^{avg}
$d \leq 20$	-	0.572	0.789	0.968	0.616	0.993	0.989	1.000	0.978	0.906	0.992	0.813	0.924	0.999	1.000
$d \leq 50$	-	0.347	0.794	0.893	0.946	0.997	0.988	1.000	0.963	0.994	0.986	0.574	0.847	0.995	1.000
All	-	<u>0.001</u>	<u>0.019</u>	0.089	0.159	0.394	0.434	0.703	0.516	0.752	0.679	<u>0.007</u>	0.062	0.437	1.000
Rank(avg)	12.59	7.19	8.57	10.34	10.79	11.82	12.81	12.8	12.52	13.50	13.34	8.60	10.63	12.44	21.43

Table 2 shows similar results for FoMo-0D w/ $D = 20$, which is pre-trained on datasets with considerably fewer dimensions. Even in this limited setting, it performs on par with the 3rd best baseline (ICL, with default HP) against 30 baselines, with an increased p -value (0.437) when compared to ICL^{avg}. On datasets with $d \leq 20$ which align with its pre-training data, it outperforms the top 5th baseline and the majority of others. With FoMo-0D pre-trained purely on synthetic datasets from a simple prior in small dimensions, these results showcase the prowess of PFNs for OD.

Figure 3 shows the distribution of AUCROC across datasets for all models (See rank distribution in Appendix Figure 17). FoMo-0D achieves a competitively small average rank with notably higher AUCROC across datasets compared to the majority of the baselines. Appendix H presents another comparison between detectors through performance profile plots (Dolan & Moré, 2002).

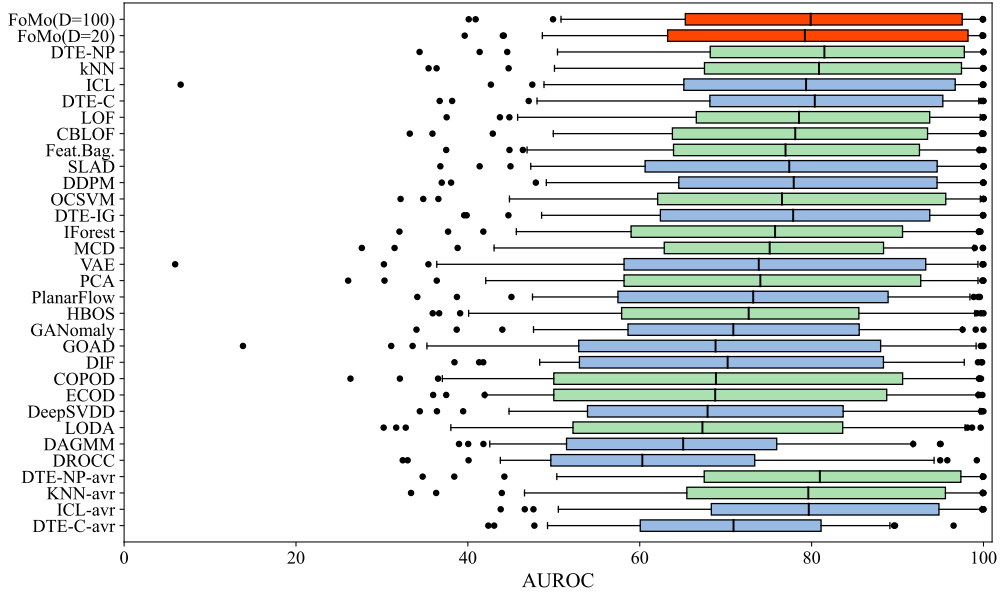


Figure 3: (best in color) AUROC (higher is better) distribution across all 57 real-world datasets shown via boxplots for (from top to bottom) FoMo-0D in red, all 26 baselines ordered by mean p -value⁶ (shallow and deep baselines in green and blue), and top 4 baselines' ^{avg} variants. The vertical line depicts the mean, the box shows the 25-75%, bars range 5-95%, and circles show the datasets at the tails.

Running time: Table 3 presents the total training time and the average inference time per test sample for FoMo-0D compared with the top-3 baselines as measured on ADBench's largest dataset. Given a new dataset, FoMo-0D bypasses model training (and HP tuning) and directly performs inference, within 7.7 ms per sample

Table 3: Train-time and Inference-time (in milliseconds) of FoMo-0D and top-3⁶ baselines (w/ *default* HPs, *excluding* time for hyperparameter optimization) on our largest dataset **donors** (see Appendix Table 20). FoMo-0D skips any model training or fine-tuning and takes a mere forward pass for inference out-of-the-box.

Method	FoMo-0D	DTE-NP	kNN	ICL
Train-time (total)	none/0-shot	1,170.29	174,157.38	130,412.86
Infer-time (per sample)	7.7	1.38	0.65	0.02

on average (see Appendix Figure 6). In comparison, all baselines need to train on each individual dataset before inference. Training time can be high for deep learning based models like ICL, further compounded by hyperparameter tuning that requires training multiple models. Even for non-parametric and/or shallow models like k NN and DTE-NP (which query k nearest neighbors), training involves various data pre-processing steps such as constructing a tree-like data structure⁷ for fast (often approximate) k NN distance querying.

4.3 Ablation Analyses

Extensive ablations in Appendix F analyze (1) the effect of D in F.1, (2) cost and performance by varying R in F.2 and F.3, (3) context size in F.4, (4) reuse periodicity P in F.5, (5) effect of data transformation T on performance and speed-up in F.6 and F.7, (6) data diversity and prolonged training in F.8, (7) quantile transformation in F.9, and (8) model size in F.10.

4.4 Generalization Analyses

Our results that synthetic pretraining at large, and GMM data prior in particular, enables FoMo-0D to reach remarkable performance on real world tasks call for deeper investigation on its generalization. To this end, Appendix G provides an extensive analysis on FoMo-0D’s generalization to out-of-distribution (OOD) synthetic datasets in G.1, GMM statistical goodness-of-fit analyses of real-world datasets in ADBench in G.2, as well as generalization to real-world OOD detection tasks in G.3.

5 Related Work

Outlier Detection (OD): Thanks to diverse applications in numerous fields, such as security, finance, manufacturing, to name a few, OD on tabular (or point-cloud) datasets has a vast literature with a long list of techniques. For earlier, shallow approaches preceding the advances in deep learning, we refer to the books by Aggarwal (2013) and Aggarwal & Sathe (2017). The modern, deep learning based techniques are surveyed in Chalapathy & Chawla (2019); Pang et al. (2021); Ruff et al. (2021). Most recent deep OD techniques take advantage of newly emerging paradigms, including self-supervised learning (Hojjati et al., 2022; Yoo et al., 2023) as well as the most recently popularized diffusion-based models (Yoon et al., 2023; Livernoche et al., 2024; Du et al., 2024; He et al., 2024).

Unsupervised Model Selection for OD: It is typical of models to exhibit various hyperparameters (HPs) that play a role in the bias-variance trade-off and hence the generalization performance, and OD models are no exception. Many earlier work on OD showed the sensitivity of classical (i.e. shallow) OD methods to the choice of their HP(s) (Aggarwal & Sathe, 2015; Campos et al., 2016; Goldstein & Uchida, 2016). Similarly, sensitivity to HPs has also been shown for deep OD models more recently (Zhao et al., 2021; Ding et al., 2022), as well as for those relying on self-supervised learning/data augmentation (Yoo et al., 2023).

While critical, work on unsupervised outlier model selection (UOMS) is slim as compared to the vast literature on detection methods. A handful of existing, mostly heuristic strategies has been studied by Ma et al. (2023) reporting discouraging results; they have shown that existing heuristics are either not significantly different from random selection, or do not outperform iForest (Liu et al., 2008) with its default HPs.

⁷Both k NN (from PyOD (Zhao et al., 2019)) and DTE-NP use BallTree from sklearn (Pedregosa et al., 2011) for nearest neighbor search; however, k NN has extra computations after calling the BallTree, which might lead to the time differences for training and inference compared to DTE-NP. We encourage the readers to refer to the official implementations for details.

Recent UOMS approaches go beyond heuristic measures and design scalable hyperensembles (Ding et al., 2022; 2024), or leverage meta-learning on historical real-world OD datasets (Zhao et al., 2021; 2022; Zhao & Akoglu, 2024). These approaches show the value of (meta)learning and transferring from historical tasks to a new task. Though grounded in the same idea of learning from a large set of (in our case, simulated) tasks, we differ in one key aspect: **FoMo-OD** is *not* a model selection technique, but rather, a foundation model that abolishes model training and selection altogether. As such, it unlocks zero(0)-shot inference on a new task.

Prior-data Fitted Networks: Based on the seminal work by Müller et al. (2022), Prior-data-fitted Networks (PFNs) establish a new paradigm for machine learning, where a PFN is pretrained on synthetic datasets generated from a data prior, and the pretrained PFN can then infer the posterior predictive distribution (PPD) for test points in a new dataset in a single forward pass, through in-context learning (Xie et al., 2021; Garg et al., 2022). It is shown that PFNs provably approximate Bayesian inference (Müller et al., 2022). Follow-up TabPFN (Hollmann et al., 2023) and its v2 Hollmann et al. (2025) achieved SOTA classification performance on small tabular datasets of size up to 1024. Other subsequent works designed LC-PFN (Adriaensen et al., 2024) and ForecastPFN (Dooley et al., 2023), respectively zero-shot learning curve extrapolation and zero-shot time-series forecasting models, trained purely on synthetic data. PFN4BO (Müller et al., 2023) employed PFNs for Bayesian optimization, while Nagler (2023) studied the statistical foundations of PFNs. Others proposed scaling the context size to enable training on larger datasets toward better generalization (Ma et al., 2024; Feuer et al., 2023; 2024; Qu et al., 2025).

Our proposed **FoMo-OD** differs from these in being the first PFN for OD, using a novel inlier/outlier data prior, employing linear transform for fast data synthesis, and incorporating the “router” attention mechanism for linear-time scalability w.r.t. context size. See Appendix K for additional details.

Zero-Shot Outlier Detection: Foundation models pretrained on massive text and image corpora, such as large language and/or vision models (L(V)LMs) like OpenAI’s GPT-series (Achiam et al., 2023), DALL-E (Ramesh et al., 2021) and Flamingo (Alayrac et al., 2022), CLIP (Radford et al., 2021), and LLaVA (Liu et al., 2024) to name a few, have demonstrated remarkable success on several zero-shot tasks in CV and NLP. Follow-up work extended these models for zero-shot out-of-distribution detection (Esmaeilpour et al., 2022), zero-shot image OD (Liznerski et al., 2022; Jeong et al., 2023; Zhou et al., 2024) as well as dialogue-based industrial image anomaly detection (Gu et al., 2024).

Foundation models, however, do not exist for tabular data which is widespread across OD applications in the real world, such as detecting credit card fraud, network intrusion, medical anomalies, and any sensor measurement abnormalities, to name a few. The recent ACR model by Li et al. (2023) on zero-shot OD does *not* rely on a pretrained foundation model, but rather is meta-trained on each specific domain using inlier-only datasets from the *same domain*. Concurrent to our work, Li et al. (2024) apply pretrained LLMs for prompt-based OD on tabular data which they serialize to text. Similar to our work, they also use *simulated* labeled OD datasets to fine-tune several existing LLMs to improve their performance. Their work, however, is quite preliminary in several fronts; a key limitation is that they assume independent features and query the LLM one-feature-at-a-time to reach an outlier score. Further, they fine-tune using only 5,000 data batches with up to 100 samples each, subsample 150 points and the first 10 columns of each dataset for evaluation (due to GPU memory constraint), and their testbed includes only two baseline methods. In contrast, **FoMo-OD** employs and pretrains PFNs at a much larger scale with rigorous evaluation on a much larger testbed.

6 Conclusion

This work introduced **FoMo-OD**, the **first foundation model for outlier detection** (OD) on tabular data. It capitalizes on the in-context learning of a Transformer model pre-trained on a large number of synthetic datasets that can then perform **zero-shot** inference on a new dataset, without *any* hyper/parameter tuning/training. **FoMo-OD** breaks new ground by fully abolishing the notoriously-hard model selection task for unsupervised OD (see Impact Statement). Further, **FoMo-OD** offers extremely fast inference thanks to a mere single forward pass. Against a long list of **26** SOTA baselines on **57** public real-world datasets, **FoMo-OD** performs on par with the *2nd* best baseline, while outperforming the majority of the baselines. Future work could expand our data prior and explore similar directions for zero-shot OD beyond tabular data. For a detailed discussion on limitations and future directions, we refer to Appendix L.

Broader Impact Statement

FoMo-0D offers zero-shot outlier detection (OD), abolishing not only parameter training but also model selection given a new dataset. This is a radical paradigm shift for the OD literature, which historically focused on designing new models and recently unsupervised model selection. Obviating the need for either, we expect FoMo-0D to route attention of the community from new model design and selection to designing better data priors and gathering datasets for pre-training, along with better and more scalable architectures for PFN.

From the applied perspective, a zero-shot OD model like FoMo-0D is a game changer for practitioners! Given the plethora of OD algorithms to choose from, each with a list of hyperparameters to set, not having the tools for effective and efficient model selection leaves the practitioners with a “choice paralysis”. With FoMo-0D, practitioners can not only bypass such dilemmas on one dataset, but thanks to the “train once, use many times” nature of pre-trained models, they can do so for any dataset, including those arriving over time. FoMo-0D is lightweight and optimized for fast inference (without any additional training or tuning), making it especially attractive for real-time or resource-constrained applications. While attractive from a performance viewpoint, we remark that FoMo-0D currently does not factor into account metrics beyond detection performance, such as fairness, biases, or other potential blindspots. In sensitive real-world applications, social and ethical costs of incorrect detections should be taken into consideration.

References

- Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altschmidt, Sam Altman, Shyamal Anadkat, et al. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*, 2023.
- Steven Adriaensen, Herilalaina Rakotoarison, Samuel Müller, and Frank Hutter. Efficient bayesian learning curve extrapolation using prior-data fitted networks. *Advances in Neural Information Processing Systems*, 36, 2024.
- Charu C. Aggarwal. *Outlier Analysis*. Springer, 2013.
- Charu C. Aggarwal and Saket Sathe. Theoretical foundations and algorithms for outlier ensembles. *Acm sigkdd explorations newsletter*, 17(1):24–47, 2015.
- Charu C. Aggarwal and Saket Sathe. *Outlier Ensembles - An Introduction*. Springer, 2017.
- Samet Akcay, Amir Atapour-Abarghouei, and Toby P Breckon. Ganomaly: Semi-supervised anomaly detection via adversarial training. In *ACCV*, pp. 622–637. Springer, 2019.
- Jean-Baptiste Alayrac, Jeff Donahue, Pauline Luc, Antoine Miech, Iain Barr, Yana Hasson, Karel Lenc, Arthur Mensch, Katherine Millican, Malcolm Reynolds, et al. Flamingo: a visual language model for few-shot learning. *Advances in neural information processing systems*, 35:23716–23736, 2022.
- Abdul Fatir Ansari, Lorenzo Stella, Caner Turkmen, Xiyuan Zhang, Pedro Mercado, Huibin Shen, Oleksandr Shchur, Syama Sundar Rangapuram, Sebastian Pineda Arango, Shubham Kapoor, et al. Chronos: Learning the language of time series. *arXiv preprint arXiv:2403.07815*, 2024.
- Lion Bergman and Yedid Hoshen. Classification-based anomaly detection for general data. In *International Conference on Learning Representations*, 2020.
- Markus M. Breunig, Hans-Peter Kriegel, Raymond T. Ng, and Jörg Sander. Lof: Identifying density-based local outliers. In *International Conference on Management of Data*, 2000.
- Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, Sandhini Agarwal, Ariel Herbert-Voss, Gretchen Krueger, Tom Henighan, Rewon Child, Aditya Ramesh, Daniel Ziegler, Jeffrey Wu, Clemens Winter, Chris Hesse, Mark Chen, Eric Sigler, Mateusz Litwin, Scott Gray, Benjamin Chess, Jack Clark, Christopher Berner, Sam McCandlish, Alec Radford, Ilya Sutskever, and Dario Amodei. Language models are few-shot learners. In *Advances in Neural Information Processing Systems*, volume 33, pp. 1877–1901, 2020.

- Guilherme O Campos, Arthur Zimek, Jörg Sander, Ricardo JGB Campello, Barbora Micenková, Erich Schubert, Ira Assent, and Michael E Houle. On the evaluation of unsupervised outlier detection: measures, datasets, and an empirical study. *Data mining and knowledge discovery*, 30:891–927, 2016.
- George Casella and Roger Berger. *Statistical inference*. CRC Press, 2024.
- Raghavendra Chalapathy and Sanjay Chawla. Deep learning for anomaly detection: A survey. *arXiv preprint arXiv:1901.03407*, 2019.
- Janez Demšar. Statistical comparisons of classifiers over multiple data sets. *The Journal of Machine learning research*, 7:1–30, 2006.
- Xueying Ding, Lingxiao Zhao, and Leman Akoglu. Hyperparameter sensitivity in deep outlier detection: Analysis and a scalable hyper-ensemble solution. *Advances in Neural Information Processing Systems*, 35: 9603–9616, 2022.
- Xueying Ding, Yue Zhao, and Leman Akoglu. Fast unsupervised deep outlier model selection with hypernetworks. *ACM SIGKDD*, 2024.
- Elizabeth D Dolan and Jorge J Moré. Benchmarking optimization software with performance profiles. *Mathematical programming*, 91:201–213, 2002.
- Samuel Dooley, Gurnoor Singh Khurana, Chirag Mohapatra, Siddhartha Venkat Naidu, and Colin White. ForecastPFN: Synthetically-trained zero-shot forecasting. In *Thirty-seventh Conference on Neural Information Processing Systems*, 2023.
- Xuefeng Du, Yiyu Sun, Jerry Zhu, and Yixuan Li. Dream the impossible: Outlier imagination with diffusion models. *Advances in Neural Information Processing Systems*, 36, 2024.
- Sepideh Esmaeilpour, Bing Liu, Eric Robertson, and Lei Shu. Zero-shot out-of-distribution detection based on the pre-trained model CLIP. In *Proceedings of the AAAI conference on artificial intelligence*, volume 36, pp. 6568–6576, 2022.
- Benjamin Feuer, Niv Cohen, and Chinmay Hegde. Scaling tabPFN: Sketching and feature selection for tabular prior-data fitted networks. In *NeurIPS 2023 Second Table Representation Learning Workshop*, 2023.
- Benjamin Feuer, Robin Tibor Schirrmester, Valeriia Cherepanova, Chinmay Hegde, Frank Hutter, Micah Goldblum, Niv Cohen, and Colin White. Tunetables: Context optimization for scalable prior-data fitted networks, 2024.
- Shivam Garg, Dimitris Tsipras, Percy S Liang, and Gregory Valiant. What can transformers learn in-context? a case study of simple function classes. *Advances in Neural Information Processing Systems*, 2022.
- Markus Goldstein and Andreas Dengel. Histogram-based outlier score (hbos): A fast unsupervised anomaly detection algorithm. *KI-2012: poster and demo track*, 1:59–63, 2012.
- Markus Goldstein and Seiichi Uchida. A comparative evaluation of unsupervised anomaly detection algorithms for multivariate data. *PloS one*, 11(4), 2016.
- Sachin Goyal, Aditi Raghunathan, Moksh Jain, Harsha Simhadri, and Prateek Jain. Drocc: Deep robust one-class classification. In *Proceedings of the 37th International Conference on Machine Learning, ICML’20*. JMLR.org, 2020.
- Zhaopeng Gu, Bingke Zhu, Guibo Zhu, Yingying Chen, Ming Tang, and Jinqiao Wang. AnomalyGPT: Detecting industrial anomalies using large vision-language models. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 38, pp. 1932–1940, 2024.
- A. Gut. *An Intermediate Course in Probability*. Springer Texts in Statistics. Springer New York, 2009. ISBN 9781441901620.

- Songqiao Han, Xiyang Hu, Hailiang Huang, Minqi Jiang, and Yue Zhao. Adbench: Anomaly detection benchmark. *Advances in Neural Information Processing Systems*, 35, 2022.
- Haoyang He, Jiangning Zhang, Hongxu Chen, Xuhai Chen, Zhishan Li, Xu Chen, Yabiao Wang, Chengjie Wang, and Lei Xie. A diffusion-based framework for multi-class anomaly detection. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 38, pp. 8472–8480, 2024.
- Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pp. 770–778, 2016.
- Jonathan Ho, Ajay Jain, and Pieter Abbeel. Denoising diffusion probabilistic models. *Advances in Neural Information Processing Systems*, 33:6840–6851, 2020.
- Hadi Hojjati, Thi Kieu Khanh Ho, and Narges Armanfard. Self-supervised anomaly detection: A survey and outlook. *arXiv preprint arXiv:2205.05173*, 2022.
- Noah Hollmann, Samuel Müller, Katharina Eggenberger, and Frank Hutter. TabPFN: A transformer that solves small tabular classification problems in a second. In *The Eleventh International Conference on Learning Representations*, 2023.
- Noah Hollmann, Samuel Müller, Lennart Purucker, Arjun Krishnakumar, Max Körfer, Shi Bin Hoo, Robin Tibor Schirrmacher, and Frank Hutter. Accurate predictions on small data with a tabular foundation model. *Nature*, 637(8045):319–326, 2025.
- Catherine Huber-Carol, Narayanaswamy Balakrishnan, Mikhail Nikulin, and Mounir Mesbah. *Goodness-of-fit tests and model validity*. Springer Science & Business Media, 2012.
- Jongheon Jeong, Yang Zou, Taewan Kim, Dongqing Zhang, Avinash Ravichandran, and Onkar Dabeer. Winclip: Zero-/few-shot anomaly classification and segmentation. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pp. 19606–19616, 2023.
- Jared Kaplan, Sam McCandlish, Tom Henighan, Tom B Brown, Benjamin Chess, Rewon Child, Scott Gray, Alec Radford, Jeffrey Wu, and Dario Amodei. Scaling laws for neural language models. *arXiv preprint arXiv:2001.08361*, 2020.
- Diederik P Kingma. Auto-encoding variational bayes. *arXiv preprint arXiv:1312.6114*, 2013.
- Diederik P. Kingma and Jimmy Ba. Adam: A method for stochastic optimization, 2017.
- Aodong Li, Chen Qiu, Marius Kloft, Padhraic Smyth, Maja Rudolph, and Stephan Mandt. Zero-shot anomaly detection via batch normalization. In *Thirty-seventh Conference on Neural Information Processing Systems*, 2023.
- Aodong Li, Yunhan Zhao, Chen Qiu, Marius Kloft, Padhraic Smyth, Maja Rudolph, and Stephan Mandt. Anomaly detection of tabular data using LLMs. *arXiv preprint arXiv:2406.16308*, 2024.
- Fei Tony Liu, Kai Ming Ting, and Zhi-Hua Zhou. Isolation forest. In *2008 Eighth IEEE International Conference on Data Mining*, pp. 413–422, 2008.
- Haotian Liu, Chunyuan Li, Qingyang Wu, and Yong Jae Lee. Visual instruction tuning. *Advances in neural information processing systems*, 36, 2024.
- Victor Liviernoche, Vineet Jain, Yashar Hezaveh, and Siamak Ravanbakhsh. On diffusion modeling for anomaly detection. In *ICLR*, 2024.
- Philipp Liznerski, Lukas Ruff, Robert A Vandermeulen, Billy Joe Franks, Klaus-Robert Müller, and Marius Kloft. Exposing outlier exposure: What can be learned from few, one, and zero outlier images. *arXiv preprint arXiv:2205.11474*, 2022.
- Junwei Ma, Valentin Thomas, Guangwei Yu, and Anthony L. Caterini. In-context data distillation with tabpfn. *CoRR*, abs/2402.06971, 2024.

- Martin Q Ma, Yue Zhao, Xiaorong Zhang, and Leman Akoglu. The need for unsupervised outlier model selection: A review and evaluation of internal evaluation strategies. *ACM SIGKDD Explorations Newsletter*, 25(1):19–35, 2023.
- Samuel Müller, Matthias Feurer, Noah Hollmann, and Frank Hutter. PFNs4BO: in-context learning for bayesian optimization. In *International Conference on Machine Learning*, 2023.
- Samuel Müller, Noah Hollmann, Sebastian Pineda-Arango, Josif Grabocka, and Frank Hutter. Transformers can do bayesian inference. *ICLR*, 2022.
- Thomas Nagler. Statistical foundations of prior-data fitted networks. In *ICML*, volume 202 of *Proceedings of Machine Learning Research*. PMLR, 2023.
- Radford M Neal. *Bayesian learning for neural networks*, volume 118. Springer Science & Business Media, 2012.
- Guansong Pang, Chunhua Shen, Longbing Cao, and Anton Van Den Hengel. Deep learning for anomaly detection: A review. *ACM computing surveys (CSUR)*, 54(2):1–38, 2021.
- Mansheej Paul, Surya Ganguli, and Gintare Karolina Dziugaite. Deep learning on a data diet: Finding important examples early in training. *Advances in neural information processing systems*, 34:20596–20607, 2021.
- Judea Pearl. *Causality*. Cambridge University Press, 2009.
- Fabian Pedregosa, Gaël Varoquaux, Alexandre Gramfort, Vincent Michel, Bertrand Thirion, Olivier Grisel, Mathieu Blondel, Peter Prettenhofer, Ron Weiss, Vincent Dubourg, et al. Scikit-learn: Machine learning in python. *the Journal of machine Learning research*, 12:2825–2830, 2011.
- Jingang Qu, David Holzmüller, Gaël Varoquaux, and Marine Le Morvan. Tabicl: A tabular foundation model for in-context learning on large data. *arXiv preprint arXiv:2502.05564*, 2025.
- Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Askell, Pamela Mishkin, Jack Clark, et al. Learning transferable visual models from natural language supervision. In *International conference on machine learning*, 2021.
- Sridhar Ramaswamy, Rajeev Rastogi, and Kyuseok Shim. Efficient algorithms for mining outliers from large data sets. In *ACM SIGMOD international conference on management of data*, 2000.
- Aditya Ramesh, Mikhail Pavlov, Gabriel Goh, Scott Gray, Chelsea Voss, Alec Radford, Mark Chen, and Ilya Sutskever. Zero-shot text-to-image generation. In *International conference on machine learning*, pp. 8821–8831. Pmlr, 2021.
- Allan Raventós, Mansheej Paul, Feng Chen, and Surya Ganguli. Pretraining task diversity and the emergence of non-bayesian in-context learning for regression. *Advances in Neural Information Processing Systems*, 2024.
- Danilo Rezende and Shakir Mohamed. Variational inference with normalizing flows. In *Proceedings of the 32nd International Conference on Machine Learning*, volume 37 of *Proceedings of Machine Learning Research*, pp. 1530–1538, Lille, France, 07–09 Jul 2015. PMLR.
- Lukas Ruff, Robert Vandermeulen, Nico Goernitz, Lucas Deecke, Shoaib Ahmed Siddiqui, Alexander Binder, Emmanuel Müller, and Marius Kloft. Deep one-class classification. In *International Conference on Machine Learning*, pp. 4393–4402, 2018.
- Lukas Ruff, Jacob R Kauffmann, Robert A Vandermeulen, Grégoire Montavon, Wojciech Samek, Marius Kloft, Thomas G Dietterich, and Klaus-Robert Müller. A unifying review of deep and shallow anomaly detection. *Proceedings of the IEEE*, 109(5):756–795, 2021.

- Tom Shenkar and Lior Wolf. Anomaly detection for tabular data with internal contrastive learning. In *International Conference on Learning Representations*, 2022.
- Ben Sorscher, Robert Geirhos, Shashank Shekhar, Surya Ganguli, and Ari Morcos. Beyond neural scaling laws: beating power law scaling via data pruning. *Advances in Neural Information Processing Systems*, 35: 19523–19536, 2022.
- Georg Steinbuss and Klemens Böhm. Benchmarking unsupervised outlier detection with realistic synthetic data. *ACM Transactions on Knowledge Discovery from Data (TKDD)*, 15(4):1–20, 2021.
- Hugo Touvron, Thibaut Lavril, Gautier Izacard, Xavier Martinet, Marie-Anne Lachaux, Timothée Lacroix, Baptiste Rozière, Naman Goyal, Eric Hambro, Faisal Azhar, Aurelien Rodriguez, Armand Joulin, Edouard Grave, and Guillaume Lample. Llama: Open and efficient foundation language models, 2023. URL <https://arxiv.org/abs/2302.13971>.
- Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N. Gomez, Lukasz Kaiser, and Illia Polosukhin. Attention is all you need. In *NIPS*, pp. 5998–6008, 2017.
- Sang Michael Xie, Aditi Raghunathan, Percy Liang, and Tengyu Ma. An explanation of in-context learning as implicit bayesian inference. *arXiv preprint arXiv:2111.02080*, 2021.
- Hongzuo Xu, Yijie Wang, Juhui Wei, Songlei Jian, Yizhou Li, and Ning Liu. Fascinating supervisory signals and where to find them: Deep anomaly detection with scale learning. In *International Conference on Machine Learning*, pp. 38655–38673. PMLR, 2023.
- Jaemin Yoo, Tiancheng Zhao, and Leman Akoglu. Data augmentation is a hyperparameter: Cherry-picked self-supervision for unsupervised anomaly detection is creating the illusion of success. *Trans. Mach. Learn. Res.*, 2023, 2023.
- Sangwoong Yoon, Young-Uk Jin, Yung-Kyun Noh, and Frank Park. Energy-based models for anomaly detection: A manifold diffusion recovery approach. *Advances in Neural Information Processing Systems*, 36:49445–49466, 2023.
- Xiaohua Zhai, Alexander Kolesnikov, Neil Houlsby, and Lucas Beyer. Scaling vision transformers. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 2022.
- Jingyang Zhang, Jingkan Yang, Pengyun Wang, Haoqi Wang, Yueqian Lin, Haoran Zhang, Yiyun Sun, Xuefeng Du, Kaiyang Zhou, Wayne Zhang, et al. Openood v1. 5: Enhanced benchmark for out-of-distribution detection. *arXiv preprint arXiv:2306.09301*, 2023.
- Yunhao Zhang and Junchi Yan. Crossformer: Transformer utilizing cross-dimension dependency for multi-variate time series forecasting. In *ICLR*, 2023.
- Yue Zhao and Leman Akoglu. Toward unsupervised outlier model selection. In *International Conference on Automated Machine Learning (AutoML)*, 2024.
- Yue Zhao, Zain Nasrullah, and Zheng Li. Pyod: A python toolbox for scalable outlier detection. *Journal of Machine Learning Research*, 20(96):1–7, 2019. URL <http://jmlr.org/papers/v20/19-011.html>.
- Yue Zhao, Ryan Rossi, and Leman Akoglu. Automatic unsupervised outlier model selection. *Advances in Neural Information Processing Systems*, 34:4489–4502, 2021.
- Yue Zhao, Sean Zhang, and Leman Akoglu. Toward unsupervised outlier model selection. In *2022 IEEE International Conference on Data Mining (ICDM)*, pp. 773–782. IEEE, 2022.
- Qihang Zhou, Guansong Pang, Yu Tian, Shibo He, and Jiming Chen. AnomalyCLIP: Object-agnostic prompt learning for zero-shot anomaly detection. In *The Twelfth International Conference on Learning Representations*, 2024. URL <https://openreview.net/forum?id=buC4E91xZE>.
- Bo Zong, Qi Song, Martin Renqiang Min, Wei Cheng, Cristian Lumezanu, Dae ki Cho, and Haifeng Chen. Deep autoencoding gaussian mixture model for unsupervised anomaly detection. In *International Conference on Learning Representations*, 2018.

Appendix

Table of Contents

We detail the contents in the appendix below.

- **Appendix A. Illustration of Synthetic Sata in 2- d** illustrates the inlier and outlier data synthesis for pre-training with a 2-dimensional example.
- **Appendix B. Linear Transform for Scalable GMM Data Synthesis** contains the proofs for efficient data synthesis in Section 3.3.
- **Appendix C. Implementation Details** includes the training and inference details of FoMo-0D.
- **Appendix D. Detailed Experiment Setup** introduces the details of pre-training and inference datasets, baselines, and their hyperparameters.
- **Appendix E. Qualitative Analysis on Sample-to-Sample Attention** visualizes the attention of FoMo-0D.
- **Appendix F. Ablation Analyses** studies different design choices of FoMo-0D.
- **Appendix G. Generalization Analyses** studies the generalization ability of FoMo-0D on out-of-distribution synthetic datasets, ADBench, and benchmarks.
- **Appendix H. Performance Profile Plots** presents a comprehensive comparison of different methods via the cumulative distribution.
- **Appendix I. Full Results** presents the detailed metric results of FoMo-0D and the baselines, including AUROC, AUPRC, and F1.
- **Appendix J. Benchmark OD Datasets** shows the details (e.g., number of samples, features) of each dataset in ADBench.
- **Appendix K. Differences to Prior Work on PFNs for Tabular Data** explains the difference and innovation of FoMo-0D from previous works.
- **Appendix L. Discussion** provides the summary of our work and discussions on the limitations and future directions of FoMo-0D.
- **Appendix M. Reproducibility Statement** details the codebase for FoMo-0D.

A Illustration of synthetic data in 2- d

We visualize our synthetic data in Figure 4, with 3 randomly created 2- d GMMs with the number of clusters ($N = 1, 2, 3$). We choose the 80th percentile as the criterion, such that inliers are samples drawn from the GMM and within the 80th percentile, and outliers are samples drawn from the inflated GMMs and outside of the 80th percentile.

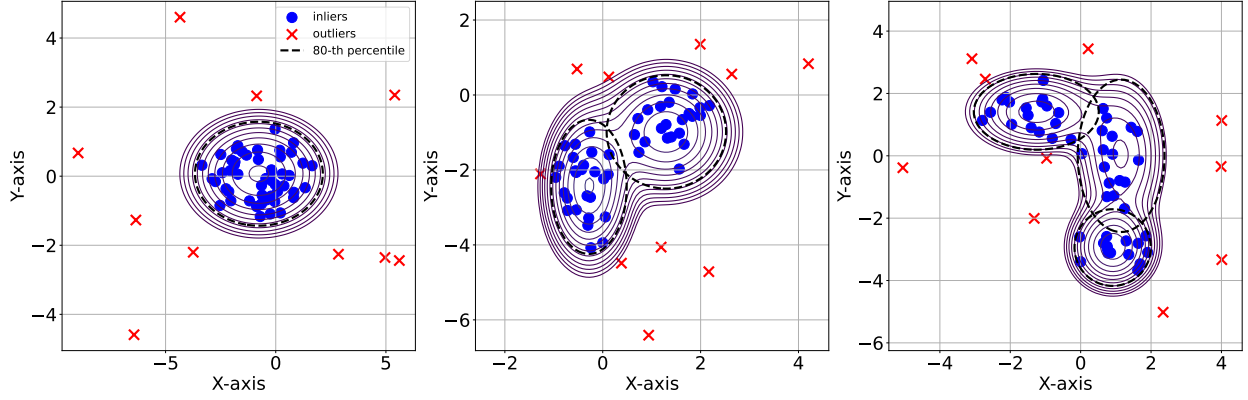


Figure 4: Illustration of synthetic data in 2D with 80th percentile as the criterion.

B Linear Transform for Scalable GMM Data Synthesis

B.1 Definitions

Definition B.1 (Gaussian Mixture Model). We denote an m -cluster d -dimension Gaussian Mixture Model as $\mathcal{G}_m^d = \{(w_j, \boldsymbol{\mu}_j, \boldsymbol{\Sigma}_j)\}_{j=1}^m$, which is the weighted sum of m Gaussian distributions:

$$p(\mathbf{x}) = \sum_{j=1}^m w_j \cdot g(\mathbf{x} | \boldsymbol{\mu}_j, \boldsymbol{\Sigma}_j), \quad (6)$$

where $w_j \in \mathbb{R}^+$ is the weight for the j -th Gaussian $\mathcal{N}(\boldsymbol{\mu}_j, \boldsymbol{\Sigma}_j)$ with $\sum_{j=1}^m w_j = 1$, and $g(\cdot | \boldsymbol{\mu}_j, \boldsymbol{\Sigma}_j)$ is the density of the j -th component/cluster, with mean/center $\boldsymbol{\mu}_j \in \mathbb{R}^d$ and covariance $\boldsymbol{\Sigma}_j \in \mathbb{R}^{d \times d}$ being positive semi-definite, such that $\mathbf{x}^T \boldsymbol{\Sigma}_j \mathbf{x} \geq 0$, for all $\mathbf{x} \in \mathbb{R}^d$.

Definition B.2 (Linear Transform). We denote a linear transformation T in $\mathbb{R}^d$ as:

$$T(\mathbf{x}) = \mathbf{W}\mathbf{x} + \mathbf{b}, \quad (7)$$

where $\mathbf{x} \in \mathbb{R}^d$, and $\mathbf{W} \in \mathbb{R}^{d \times d}$, $\mathbf{b} \in \mathbb{R}^d$ are the parameters of T .

Definition B.3 (Mahalanobis Distance). The Mahalanobis distance dist_M between a point $\mathbf{x} \in \mathbb{R}^d$ and a Gaussian distribution $\mathcal{N}(\boldsymbol{\mu}, \boldsymbol{\Sigma})$ is defined as:

$$\text{dist}_M(\mathbf{x}) = \sqrt{(\mathbf{x} - \boldsymbol{\mu})^T \boldsymbol{\Sigma}^{-1} (\mathbf{x} - \boldsymbol{\mu})}. \quad (8)$$

Definition B.4 (χ_d^2 -distribution). The Chi-squared distribution χ_d^2 with d degrees of freedom is the distribution of the sum of squares of d independent standard Normal random variables.

B.2 Properties

Property B.5 (Lemma 5.3.2 (Casella & Berger, 2024)). If $Z \sim \mathcal{N}(0, 1)$, then $Z^2 \sim \chi_1^2$; If $X_1, \dots, X_d$ are independent and $X_i \sim \chi_1^2$, then $\sum_{i=1}^d X_i \sim \chi_d^2$.

Property B.6. The squared Mahalanobis distance $\text{dist}_M^2(\mathbf{x}) \sim \chi_d^2$, with $\mathbf{x} \sim \mathcal{N}(\boldsymbol{\mu}, \boldsymbol{\Sigma})$.

Proof: If $\mathbf{x} \sim \mathcal{N}(\boldsymbol{\mu}, \boldsymbol{\Sigma})$, then we have $\mathbf{z} = \boldsymbol{\Sigma}^{-\frac{1}{2}}(\mathbf{x} - \boldsymbol{\mu}) \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_d)$ (Gut, 2009), such that:

$$\text{dist}_M^2(\mathbf{x}) = \mathbf{z}^T \mathbf{z} = \sum_{i=1}^d z_i^2 \quad (9)$$

where z_i are independent standard Normal random variables. We have $\sum_{i=1}^d z_i^2 \sim \chi_d^2$ from Property B.5, which completes the proof.

B.3 Lemmas

Lemma B.7. Linear transform T with invertible $\mathbf{W}$ on $\mathcal{G}_m^d$ preserves Mahalanobis distances.

Proof: We denote the transformed GMM as $T(\mathcal{G}_m^d) = \{(w_j, \mathbf{W}\boldsymbol{\mu}_j + \mathbf{b}, \mathbf{W}\boldsymbol{\Sigma}_j\mathbf{W}^T)\}_{j=1}^m$, then with $\mathbf{x} \sim \mathcal{N}(\boldsymbol{\mu}_j, \boldsymbol{\Sigma}_j)$, for the transformed point $T(\mathbf{x})$ we have:

$$\text{dist}_M(T(\mathbf{x})) = \sqrt{(T(\mathbf{x}) - (\mathbf{W}\boldsymbol{\mu}_j + \mathbf{b}))^T (\mathbf{W}\boldsymbol{\Sigma}\mathbf{W}^T)^{-1} (T(\mathbf{x}) - (\mathbf{W}\boldsymbol{\mu}_j + \mathbf{b}))} \quad (10)$$

$$= \sqrt{(\mathbf{W}(\mathbf{x} - \boldsymbol{\mu}_j))^T (\mathbf{W}\boldsymbol{\Sigma}\mathbf{W}^T)^{-1} (\mathbf{W}(\mathbf{x} - \boldsymbol{\mu}_j))} \quad (11)$$

$$= \sqrt{(\mathbf{x} - \boldsymbol{\mu}_j)^T \mathbf{W}^T (\mathbf{W}^T)^{-1} \boldsymbol{\Sigma}^{-1} \mathbf{W}^{-1} \mathbf{W} (\mathbf{x} - \boldsymbol{\mu}_j)} \quad (12)$$

$$= \sqrt{(\mathbf{x} - \boldsymbol{\mu}_j)^T \boldsymbol{\Sigma}^{-1} (\mathbf{x} - \boldsymbol{\mu}_j)} = \text{dist}_M(\mathbf{x}) . \quad (13)$$

□

Lemma B.8. Linear transform T with invertible $\mathbf{W}$ on $\mathcal{G}_m^d$ preserves the percentiles of the GMM.

Proof: Let $\chi_d^2(\alpha)$ denote the α -th percentile of χ_d^2 , such that for $X \sim \chi_d^2$:

$$\text{Prob}(X \leq \chi_d^2(\alpha)) = \frac{\alpha}{100} . \quad (14)$$

Based on Property B.6, we have $\text{Prob}(\text{dist}_M^2(\mathbf{x}) \leq \chi_d^2(\alpha)) = \frac{\alpha}{100}$.

Let $\mathbf{x} \sim \mathcal{G}_m^d$, such that $\text{dist}_M^2(\mathbf{x}) > \chi_d^2(\alpha)$ for all $\mathcal{N}_j(\boldsymbol{\mu}_j, \boldsymbol{\Sigma}_j)$, which indicates that $\mathbf{x}$ is outside the α -th percentile of $\mathcal{G}_m^d$. Since $\text{dist}_M(\mathbf{x})$ is preserved under T (see Lemma B.7), then we conclude that the linear transform T with invertible $\mathbf{W}$ preserves the percentiles of the GMM. □

C Implementation details

C.1 Hardware

We base our experiments on a NVIDIA RTX A6000 GPU with AMD EPYC 7742 64-Core Processors.

C.2 Training and inference

We train our models for 200 epochs with the Adam optimizer (Kingma & Ba, 2017) and a `learning_rate` = 0.001, and test with the model corresponding to the lowest training loss. The size of our $D = \{20, 100\}$ model is 4.87M and 4.89M parameters, respectively. We show the training process of PFNs and our model in Algorithm 1.

Dealing with varying dimensions and dataset size To handle input with varying number of d features, we follow Müller et al. (2022). Specifically for $d < D$, we rescale the input with $\frac{D}{d}$ and pad the features to size D with 0s; and for $d > D$, we randomly sample D features out of d . In addition, FoMo-0D uses context size of up to 5K at inference, where for each test sample $\mathbf{x} \in \mathcal{D}_{\text{test}}$, we randomly sample $(5K-1)$ points as $\mathcal{D}_{\text{train}}$ from datasets with $n > 5K$.

Algorithm 1: Prior-fitting of a PFN (Müller et al., 2022) and **ours**

Input : A prior distribution over datasets $p(\mathcal{D})$, from which samples can be drawn and the number of datasets Q to draw for one epoch, the number of training epochs E , the periodicity P , the number of unique datasets q , linear transformation T .

Output : A model q_θ that will approximate the PPD

```

1 Initialize the neural network  $q_\theta$ ;
2 Initialize the epoch-level collection  $\mathcal{C}_E = []$ ;
3 for  $i \leftarrow 1$  to  $E$  do
4   if  $i \leq P$  then
5     Initialize an empty buffer  $\mathcal{B}_i = []$ ;
6     Initialize the dataset-level collection  $\mathcal{C}_q = []$ ;
7     for  $j \leftarrow 1$  to  $Q$  do
8       if  $j \leq q$  then
9         Step 1: sample  $D_j := \mathcal{D}_{\text{train}} \cup \{(\mathbf{x}_k, y_k)\}_{i=k}^{|\mathcal{D}_{\text{test}}|} \sim p(\mathcal{D})$ ;
10         $\mathcal{C}_q \leftarrow \mathcal{C}_q + [D_j]$ 
11      end
12    else
13       $j \leftarrow j \bmod q$ 
14       $D_j \leftarrow T(\mathcal{C}_q[j])$ 
15    end
16    Step 2: compute stochastic loss approximation  $\bar{\ell}_\theta = \sum_{k=1}^{|\mathcal{D}_{\text{test}}|} (-\log q_\theta(y_k | \mathbf{x}_k, \mathcal{D}_{\text{train}}))$ ;
17    Step 3: update parameters  $\theta$  with stochastic gradient descent on  $\nabla_\theta \bar{\ell}_\theta$ ;
18     $\mathcal{B}_i \leftarrow \mathcal{B}_i + [D_j]$ 
19  end
20   $\mathcal{C}_E \leftarrow \mathcal{C}_E + [\mathcal{B}_i]$ 
21 end
22 else
23    $i \leftarrow i \bmod P$ 
24    $\mathcal{B}_i \leftarrow \mathcal{C}_E[i]$ 
25   for  $j \leftarrow 1$  to  $Q$  do
26      $D_j \leftarrow T(\mathcal{B}_i[j])$ 
27     Perform Step 2 and Step 3
28   end
29 end
30 end

```

Model architecture We use a 4-layer Transformer with hidden dimension $\mathbf{h_dim} = 256$, a linear embedding layer at the input ($\mathbb{R}^D \rightarrow \mathbb{R}^{\mathbf{h_dim}}$), and a 2-layer MLP layer at the output ($\mathbb{R}^{\mathbf{h_dim}} \rightarrow \mathbb{R}^2$) for inlier vs. outlier binary classification. For each Transformer layer, we use $\mathbf{num_head} = 4$ for each attention module and $R = 500$ for the router-based attention (Figure 2).

Training loss In Figure 5, we plot the training loss of our $D = 100$ model trained with 8K unique datasets/epoch (denoted as “8K”) versus 0.5K unique + 7.5K transformed datasets/epoch (denoted as “0.5K+T”), together with the $D = 20$ model trained with reuse periodicity $P = 1$ (denoted as “P=1”, reusing the same 8K datasets across epochs) and $P = 1$ with transformation (denoted as “P=1+T”, transforming the 8K datasets across epochs). Notice that the loss with transformation is slightly higher than no transformation (i.e., $D = 100$, “0.5K+T” vs. “8K”, and $D = 20$, “P=1+T” vs. “P=1”) across all 200 epochs, which is reasonable since the transformed datasets have non-diagonal covariances that make the learning task harder and thus result in a higher training loss. The training losses of FoMo-0D with $D = 100$ are also higher than with $D = 20$ since the subspace OD tasks are harder in higher dimensions.

Inference time Figure 10 (left) showed the inference time of FoMo-0D on CPU, comparing typical attention versus the router-based attention (with $R = 500$ routers) under varying context sizes from 1K to 10K. The

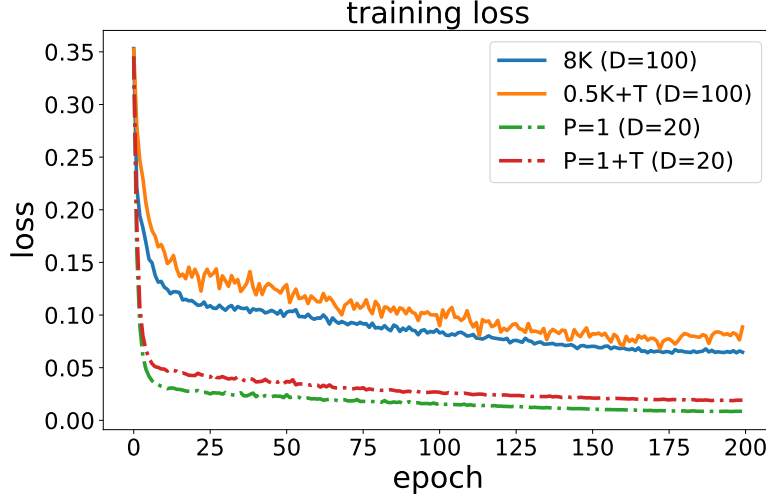


Figure 5: (best in color) Training loss of FoMo-0D ($D = 100$) with 8K unique datasets/epoch (in blue) and using 0.5K unique + 7.5K transformed datasets/epoch (in orange), and FoMo-0D ($D = 20$) with $P = 1$ (in green) and $P = 1$ with transformation (in red) over 200 epochs.

time is measured on CPU to clearly showcase the scalability trends; *quadratic* without routers and *linear* with routers.

Figure 6 shows the inference time on GPU. Notice that the time is much lower (in milliseconds), thanks to the Transformer architecture taking advantage of GPU parallelism, while the compute time for attention without routers continues to grow faster than that with routers.

In implementation, FoMo-0D (with $R = 500$ routers) uses inference context size of 5K by default, which takes about 7.7 ms per test sample on average.

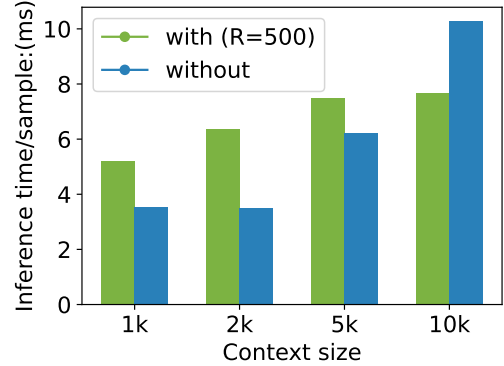


Figure 6: Inference time of FoMo-0D on GPU with vs. w/out router-based attention under varying context size.

D Detailed Experiment Setup

D.1 Pre-training Dataset Synthesis

During pretraining, we generate unique GMM datasets by first drawing a configuration, including dimensionality $d \in [D]$, number of components $m \in [M]$, centers $\{\mu_j\}_{j=1}^m$ (each $\mu_j \in [-5, 5]^d$) and covariances $\{\Sigma_j\}_{j=1}^m$ ($\text{diag}(\Sigma_j) \in [-5, 5]^d$). We set $M = 5$ and vary $D \in \{20, 100\}$ to study pretraining with relatively small and high dimensional datasets, respectively. We synthesize inliers and outliers as described in Section 3.1.

We then sample $S = 5,000$ points that are within the 90th percentile of the GMM. To synthesize outliers, we “inflate” a *subset* of dimensions by randomly choosing $|\mathcal{K}| \in [D]$ dimensions and multiplying the corresponding variances by $\times 5$ (following (Han et al., 2022)), i.e. $5 \times \Sigma_{j,kk}$ ’s for $k \in \mathcal{K}$, and then draw $S = 5,000$ samples from the inflated GMM that are outside the 90th percentile of the original GMM.

To speed up data synthesis via linear transformations, we first draw 500 unique datasets using $m \in [5]$ and $d \in \{1, 2, \dots, 100\}$ (i.e. 5×100) and transform each one $15 \times$ using varying parameters $(\mathbf{W}, \mathbf{b})$ as described in Section 3.3.⁸ This yields 8K unique datasets (500 original and 7,500 transformed) to use at one training epoch (over 1,000 steps with batch size $B = 8$). We repeat this process at each epoch, drawing 500 new datasets and transforming them to reach 8K datasets per epoch.

⁸It is important to ensure that the eigenvalues of $\mathbf{W}$ (i.e. variances) are not too small such that the dataset does not flatten in any direction. To this end, we draw a random orthonormal basis $\mathbf{U} \in [-1, 1]^{d \times d}$ and a diagonal $\mathbf{\Lambda}$ with eigenvalues $\lambda_{kk} \in ([-1, -0.1] \cup [0.1, 1])^d$, and obtain $\mathbf{W} = \mathbf{U}\mathbf{\Lambda}\mathbf{U}^T$. We also use $\mathbf{b} \in [-1, 1]^d$.

D.2 Real-world Benchmark Datasets

While pretraining is purely on synthetic datasets, we evaluate FoMo-OD on **57** real-world datasets from the ADBench benchmark (Han et al., 2022) (see Table 20). They consist of 47 popular tabular outlier detection datasets, as well as 10 newly-constructed tabular datasets created from images and natural language tasks by using pretrained models to extract embeddings. We defer to the original paper for the details on these benchmark datasets.

We compare to DTE (Livernoche et al., 2024) and baselines therein as described next, thus, following their OD setting with inlier-only $\mathcal{D}_{\text{train}}$, we split each dataset five times into train/test using five different seeds and report the mean performance and its standard deviation. In particular, each random split designates 50% of the inliers as $\mathcal{D}_{\text{train}}$, while $\mathcal{D}_{\text{test}}$ contains the rest of the inliers and all the outlier samples. Note that while the baseline methods require model re-training and inference for each $\mathcal{D}_{\text{train}}/\mathcal{D}_{\text{test}}$ split, FoMo-OD uses the splits only for inference as $\mathcal{D}_{\text{train}}$ is merely passed as context.

Before passing the datasets as input to FoMo-OD, we perform a quantile transform such that the features follow a Normal distribution, to better align with the pretraining data from GMMs.

D.3 Baselines

We compare FoMo-OD against **26** baselines, from classical/shallow methods to modern/deep models. Our baselines include all the baselines imported from one of the latest papers that proposed the SOTA diffusion-based model DTE (Livernoche et al., 2024), and its three variants; DTE-C, DTE-IG, and DTE-NP. Their baselines comprise all those in ADBench (Han et al., 2022); both classical ones (k NN (Ramaswamy et al., 2000), LOF (Breunig et al., 2000), iForest (Liu et al., 2008), HBOS (Goldstein & Dengel, 2012), etc.) and deep models (DeepSVDD (Ruff et al., 2018), DAGMM (Zong et al., 2018), DROCC (Goyal et al., 2020), etc.). They also include more recent approaches based on self-supervised learning (GOAD (Bergman & Hoshen, 2020), ICL (Shenkar & Wolf, 2022), SLAD (Xu et al., 2023), etc.), besides the four additional generative baselines: normalizing planar flows (Rezende & Mohamed, 2015), DDPM (Ho et al., 2020), VAE (Kingma, 2013) and GANomaly (Akçay et al., 2019). We defer to the original paper for additional details. Overall, our 26 baselines consist of the most recent, SOTA approaches for OD that span a diverse family (nonparametric, self-supervised, generative, etc.).

D.4 Hyperparameters for Baselines

Table 4 gives the list of HP values we used to study the HP sensitivity/performance variability of the (from top to bottom) top-4 baselines.

Table 4: Top-4 baselines (from top to bottom) and hyperparameter (HP) configurations.

Baseline	Hyperparameters
DTE-NP	$k \in \{5, 10, 20, 40, 50\}$
k NN	$k \in \{5, 10, 20, 40, 50\}$
ICL	<code>learning_rate</code> $\in \{10^{-1}, 10^{-2}, 10^{-3}, 10^{-4}, 10^{-5}\}$
DTE-C	$k \in \{5, 10, 20, 40, 50\}$

D.5 Ranking the 26 baselines

Figure 24 presents the visualization of the p -values of the pairwise Wilcoxon signed rank test w.r.t. AUROC among the baseline methods used by Livernoche et al. (2024). We rank these 26 baselines based on their mean p -value (i.e., row-wise average) against the other baselines.

D.6 Comparison of top-4 baseline variants with varying HP configurations

Figure 25, 26, 27, 28 give the p -values, respectively comparing the variants of the top-4 baselines (DTE-NP, k NN, ICL, DTE-C) among themselves using different HP configurations, as well as the avg model with the average performance across HPs. (Specifically for ICL, `learning_rate` ($1r$) $\in \{10^{-1}, 10^{-2}, 10^{-3}, 10^{-4}, 10^{-5}\}$; and for others, `#nearest-neighbors` $k \in \{5, 10, 20, 40, 50\}$). We find that for ICL, $1r = 10^{-3}$ or 10^{-4} are preferable while those that are too small or too large perform poorly. For others, small $k \in \{5, 10\}$ tend to outperform larger $k \in \{40, 50\}$. Note that Livernoche et al. (2024) used $k = 5$ in their paper that proposed DTE (and variants) as well as the k NN baseline for fair comparison, while the DTE^{avg} and $k\text{NN}^{\text{avg}}$ models across HP configurations perform subpar.

D.7 Sampling time of d -dimensional GMM

Figure 7 shows the sampling time of drawing 10,000 points from different GMMs with increasing dimensionality $d = \{10, 20, \dots, 200\}$. We parallelize the sampling process over 10 CPUs, where each CPU draws 1000 samples.

We observe that the sampling time grows nonlinearly as the number of dimensions increases, which suggests that it may incur considerable computational overhead to directly draw from the data prior over hundreds of thousands of training steps, motivating the use of our proposed on-the-fly linear transformation T for scalability.

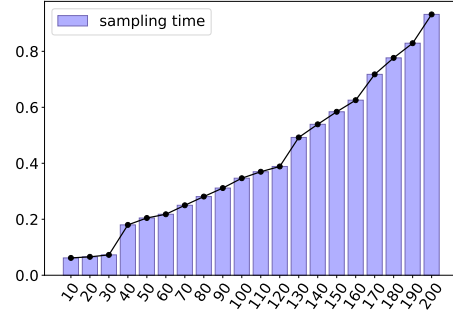


Figure 7: Sampling time (in seconds) of 10,000 points from GMMs with varying number of dimensions.

E Qualitative Analysis on Sample-to-Sample Attention

We sample 50 inliers as context and 100 outliers from a 2- d GMM using the 80th percentile as the labeling threshold, and visualize the top 5 inliers most attended by the 100 outliers based on the average (cross) attention weights over 4 heads from the last layer of FoMo-0D ($D = 100$), which accurately labeled all the 100 outliers. In Figure 8, the most frequently attended inliers are close to either the center of a Gaussian (e.g., 1st, 5th) or the criterion (e.g., 3rd, 4th), suggesting FoMo-0D tends to learn decision boundaries that reflect the prior data generation process.

For each outlier, we compute the sum of L2 distances to its top-5 attended inliers (`att`), the sum of L2 distances to 5 randomly chosen inliers (`rdm`), and the sum of L2 distances to top-5 inliers with highest likelihood under the GMM (`prob`). We perform Wilcoxon signed rank test between `att` and `rdm` (alternative: “less”), `att` and `prob` (alternative: “greater”) over all the outliers, with a p -value of 4.4×10^{-4} and 0.99, respectively, suggesting the distances based on attention weights are significantly less than the random distances, and **not** significantly greater than the distances to inliers in high probability region.

We visualize the top-5 attended inliers for 3 outliers at different position of the 2- d GMM in Figure 9. For a specific outlier, there is a similar trend of attending to the center of a Gaussian (as shown in Figure 8), besides, inliers that reflect the criterion boundary or are close to the outlier are actively attended (e.g., 3rd, 4th in the left, 1st in the middle, 2nd, 5th in the right), suggesting FoMo-0D is incorporating both boundary and nearest neighbor information dynamically for each outlier.

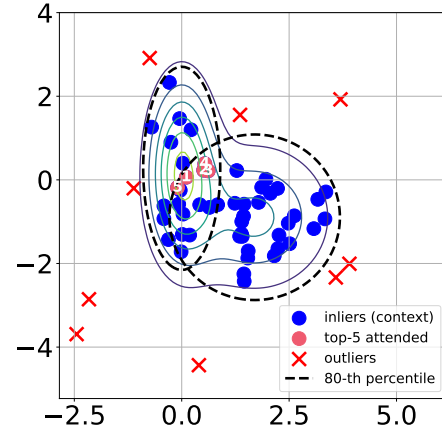


Figure 8: Top-5 attended inliers (all 50 inliers and only part of the outliers are shown for better visualization).

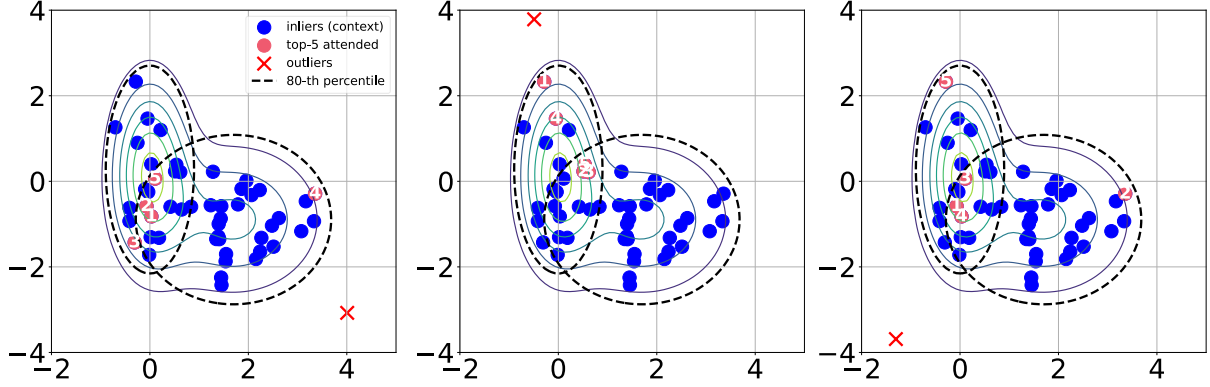


Figure 9: Top-5 attended inliers of 3 outliers at different positions of the GMM

F Ablation Analyses

In this section, we perform various ablations to study the effect of different design choices in FoMo-0D; namely, **F.1** maximum pretraining data dimensionality D , the number of routers R on **F.2** cost and **F.3** performance, **F.4** context size (both for training and inference), **F.5** number of unique datasets used for pretraining (i.e., reuse periodicity P), data transformation T during synthesis on **F.6** performance and **F.7** speed up, **F.8** data diversity and prolonged training, **F.9** quantile transforming the benchmark datasets preceding inference, and finally, **F.10** how different model sizes affect performance.

Unless stated otherwise, most ablation results are performed using FoMo-0D with $D = 20$, as it is faster to pretrain under these many varying settings.

F.1 Effect of pretraining dimensionality D

How does FoMo-0D’s generalization performance change by increasing dimensionality of the pretraining data?

We start by comparing FoMo-0D pretrained on datasets with up to $D = 20$ versus $D = 100$ dimensions. Note that learning on higher dimensional datasets is harder, as evident from the relatively larger pretraining loss as shown in Appendix Figure 5. While the statement is accurate in general, it is also partly because subspace outliers “hide” better in higher dimensions.

Comparing Table 1 ($D = 100$) with Table 2 ($D = 20$) w.r.t. p -values over All datasets, we find that FoMo-0D at larger scale does better, where **all** p -values are larger for $D = 100$ than $D = 20$. We find that FoMo-0D with $D = 20$ performs well on datasets with $d \leq 20$ (i.e., “on its own game”), however beyond its pretraining setting, e.g. on datasets with $d \leq 50$, $D = 100$ is superior to $D = 20$ as shown in Appendix Table 13.

F.2 Effect of routers on cost

What is the running time and memory cost of FoMo-0D with $\mathcal{E}$ w/out router-based attention?

Figure 10(left) shows the average inference time per test sample, comparing FoMo-0D using a router-based attention mechanism with $R = 500$ routers (in green) versus FoMo-0D using typical attention without any routers (in blue). As inference context size increases, running time for traditional attention grows quadratically while router mechanism scales linearly.⁹

Similarly, memory cost with routers is considerably lower when using routers, especially for larger context sizes, as shown in Figure 10(middle).

⁹Note that the inference time is reported on CPUs to show scalability. On GPUs, w/ 5K context size, see Appendix Figure 6, where typical attention takes advantage of parallelism (6.5ms), while router-based attention is slightly slower (7.7 ms w/ 500 routers) due to its **two** sequential self-attentions; see Eq.s (4) and (5).

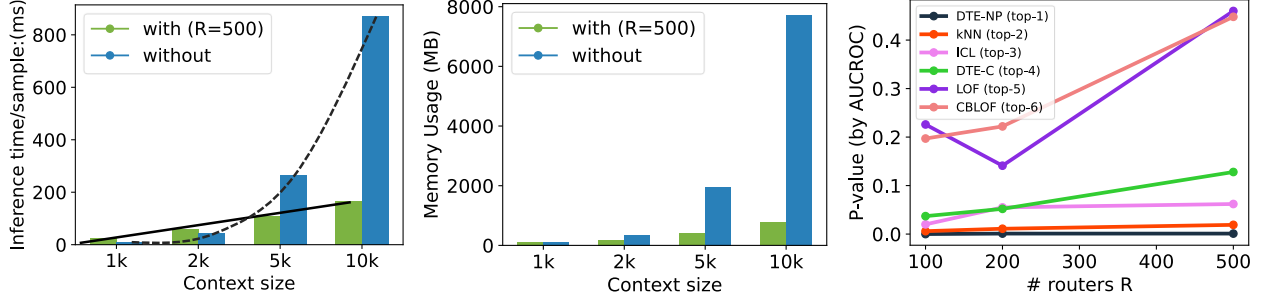


Figure 10: FoMo-0D w/ router mechanism saves time and memory while more #routers perform better, offering a cost-performance trade-off: (left) inference-time (ms) per sample and (middle) memory cost (MB) with & w/out routers by varying context size; (right) performance (based on p -value against top baselines, higher is better) vs. number of routers. (setting: $D = 20$, $P = 1$)

F.3 Effect of routers on performance

What is the impact of the number R of routers (or representatives) on performance?

Router-based mechanism allows to trade-off running time with expressiveness of the attention and hence performance. Figure 10(right) shows the p -values of the Wilcoxon signed rank test as the number of routers R is increased from 100 to 200 and 500, comparing FoMo-0D to each of the top-6 baselines. We notice that FoMo-0D performance tends to increase monotonically with more routers.

F.4 Effect of context size

What is the impact of context size, both during model pretraining as well as during inference?

To study how performance changes by context size, we train FoMo-0D with varying context size in $\{1K, 2K, 5K\}$ and employ each pretrained model for inference with varying context size in $\{1K, 2K, 5K, 10K\}$. Table 5 shows the results, where performance is depicted by the average rank of FoMo-0D (the lower, the better).

Table 5: Average rank (based on comparison to 30 baselines w.r.t. AUROC) of FoMo-0D across datasets under *different context sizes* for training and inference. Smaller ranks imply better performance. (setting: $D = 20$, $R = 500$, $P = 1$)

	Infer:1K	Infer:2K	Infer:5K	Infer:10K
Train:1K	13.816	14.623	15.193	15.439
Train:2K	13.079	13.219	13.439	13.561
Train:5K	13.088	13.211	13.307	13.430

We find that training with a larger context improves performance at any inference context size. On the other hand, perhaps counter-intuitively, FoMo-0D with smaller inference context size does better. We conjecture that is because the #routers-to-context size ratio increases with a larger context size at inference, limiting the expressive power of the “bottleneck” attention mechanism. The pairwise statistical tests among the $3 \times 4 = 12$ models support these observations, as shown in Figure 11. Interestingly, when the training context size is large enough at 5K, inference with 10K samples generalizes beyond training with no statistical evidence for performance difference (at 0.05) from other inference context sizes.

F.5 Effect of number of unique datasets

How do FoMo-0D performances compare when pretrained on unique vs. reused datasets, via varying periodicity P ?

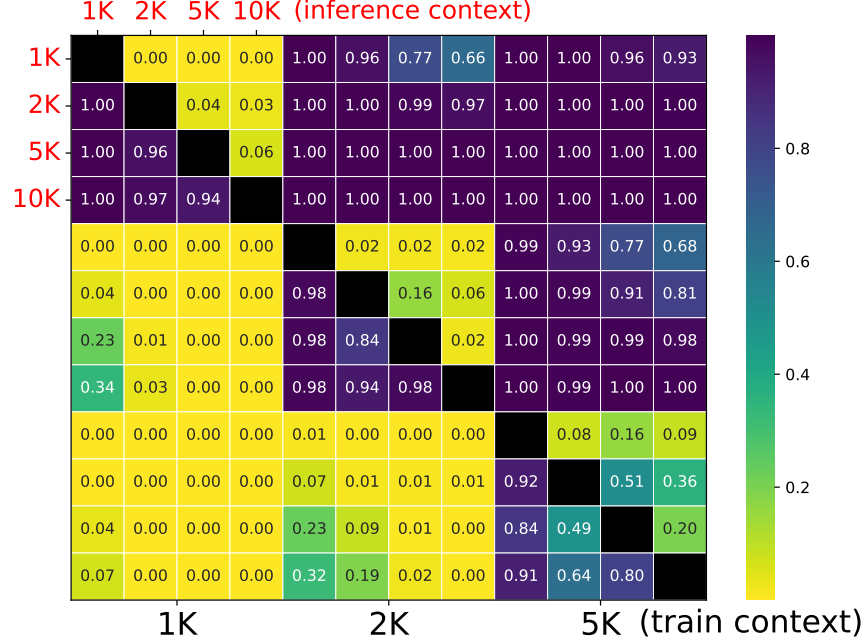


Figure 11: p -values of the pairwise Wilcoxon signed rank test between models (larger p implies col-method is better than row-method) w/ different context sizes for **training** (1K/2K/5K, 1st/2nd/3rd four grids, in **black**) and **inference** (1K/2K/5K/10K, every 1st/2nd/3rd/4th grid, in **red**): Larger training context improves overall performance, while smaller inference context is preferable.

Table 6: Ablation results on dataset reuse across epochs with varying $P \in \{1, 50, 100\}$ show stable p -values against the top-5 baselines, where there is no statistical evidence to suggest performance difference between FoMo-0D with $D = 20$ and the top 3rd baseline at 0.05 w.r.t. pairwise Wilcoxon signed rank test comparisons, while it continues to significantly outperform the top 5th baseline (LOF) when $d \leq 50$. (setting: $D=20$, $R=500$, context size=5K, w/out transformation T)

	$P = 1$ (#unique datasets: 8K)					$P = 50$ (#unique datasets: $8 \times 50 = 400K$)					$P = 100$ (#unique datasets: $8 \times 100 = 800K$)				
top-5	DTE-NP	kNN	ICL	DTE-C	LOF	DTE-NP	kNN	ICL	DTE-C	LOF	DTE-NP	kNN	ICL	DTE-C	LOF
All	0.001	0.019	0.062	0.128	0.460	0.001	0.019	0.089	0.159	0.394	0.001	0.015	0.072	0.121	0.290
$d \leq 20$	0.583	0.755	0.943	0.736	0.998	0.572	0.789	0.968	0.616	0.993	0.439	0.678	0.953	0.550	0.972
$d \leq 50$	0.415	0.750	0.869	0.962	0.999	0.347	0.794	0.893	0.946	0.997	0.293	0.697	0.890	0.924	0.994

Next we study the effect of dataset *reuse at epoch level* (w/out transformation) on performance as presented in Section 3.3. We vary reuse periodicity P in $\{1, 50, 100\}$, and accordingly, increase the number of unique datasets used for pretraining across epochs. As shown in Table 6, FoMo-0D (w/ $D = 20$) performs similarly with varying dataset reuse. In fact, it is competitive even with $P = 1$, remaining no different from the 3rd best baseline (ICL) across All (57) datasets, while significantly outperforming the top 5th (LOF) across (24) datasets with $d \leq 20$ as well as (38) with $d \leq 50$.

F.6 Effect of transformation T for synthesis

How do FoMo-0D performances compare when pretrained on datasets with vs. w/out linear transformation?

Setting $P = 1$, we next study the impact of linear transformation T . Table 7 presents the results, where we compare reuse of the *same* 8K unique datasets across epochs (w/out T), versus *transforming* these datasets with T at every epoch with different parameters (w/ T). FoMo-0D performance remains stable; no statistical evidence for performance difference from the top 3rd model on All datasets, while significantly outperforming the top 5th across those with $d \leq 20$ and $d \leq 50$. This suggests that T can be employed without sacrificing performance to save time during pretraining.

Table 7: Ablation results on performance w/ & w/out linear transformation T show stable p -values against the top-5 baselines, with no statistical evidence for performance difference between FoMo-0D with $D = 20$ and the top 3rd baseline at 0.05 w.r.t. pairwise Wilcoxon signed rank test comparisons. (setting: $D = 20$, $R = 500$, context size=5K, $P = 1$)

	w/out transformation T					w/ transformation T				
top-5	DTE-NP	kNN	ICL	DTE-C	LOF	DTE-NP	kNN	ICL	DTE-C	LOF
All	<u>0.001</u>	<u>0.019</u>	0.062	0.128	0.460	<u>0.002</u>	<u>0.015</u>	0.226	0.210	0.280
$d \leq 20$	0.583	0.755	0.943	0.736	0.998	0.648	0.708	0.988	0.718	0.955
$d \leq 50$	0.415	0.750	0.869	0.962	0.999	0.264	0.382	0.971	0.900	0.963

F.7 Speed up by T

What is the time saving on data synthesis with linear transformation?

Figure 12 shows the distribution of pretraining running-time per epoch with and w/out data transformation. Specifically, we compare (left) generating 8K unique datasets/epoch on-the-fly and (right) first generating 500 unique datasets on-the-fly and then transforming each one 15 times using T with different parameters to reach 8K datasets at each epoch.

Notice that pretraining with T takes about 450 sec./epoch on average, while without T it requires 1200 sec./epoch to generate 8K unique datasets and gradient descent across 1000 steps. Different from other ablation results, which are based on the $D = 20$ model, here we report the running times for our $D = 100$ model. Overall, our final FoMo-0D took ≈ 25 hours for pre-training (450 sec. $\times$ 200 epochs). Importantly, this is a one-time cost that amortizes across many downstream tasks with as low as **7.7 ms inference time** per test sample (see Table 3 and Appendix Figure 6).

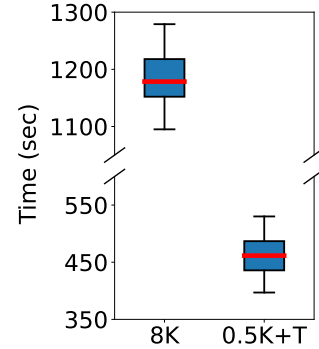


Figure 12: Runtime/epoch distribution over 100 epochs for FoMo-0D ($D=100$) with (left) $P=100$, i.e. 8K unique datasets/epoch vs. (right) 0.5K unique+7.5K transformed datasets/epoch.

F.8 Effect of data diversity and prolonged training

How does FoMo-0D's performance change by increasing pretraining data diversity and number of training epochs?

Originally we have trained FoMo-0D w/ $D = 100$ using 0.5K unique + 7.5K transformed datasets over 200 epochs. As mentioned earlier, learning in higher dimensions tends to incur a larger loss in general but also specifically here, as subspace outliers are harder to detect in high dimensions.

Toward reducing the loss further, we resume the pretraining for another 100 epochs. Further, to simplify the tasks and thereby increase data diversity, we also decrease the inlier/outlier labeling percentile threshold from 90% to 80% during on-the-fly data generation in the last 100 epochs. In Figure 13, we present the training loss of FoMo-0D ($D = 100$) trained with 0.5K unique + 7.5K transformed datasets/epoch over 200 epochs (90th percentile as labeling threshold) and then 100 additional epochs (80th percentile as the threshold) to show how data diversity and amount affect model performance.

Figure 14 compares FoMo-0D's performance (w/ $D = 100$) to top-5 baselines w.r.t. p -values of the paired Wilcoxon signed rank test on datasets with $d \leq 100$, after the first 200 epochs versus after 300 epochs. The increase in all the p -values showcases the benefit of additional training.

F.9 Effect of applying quantile transform on benchmark datasets

What is the impact of quantile data transform preceding inference on performance?

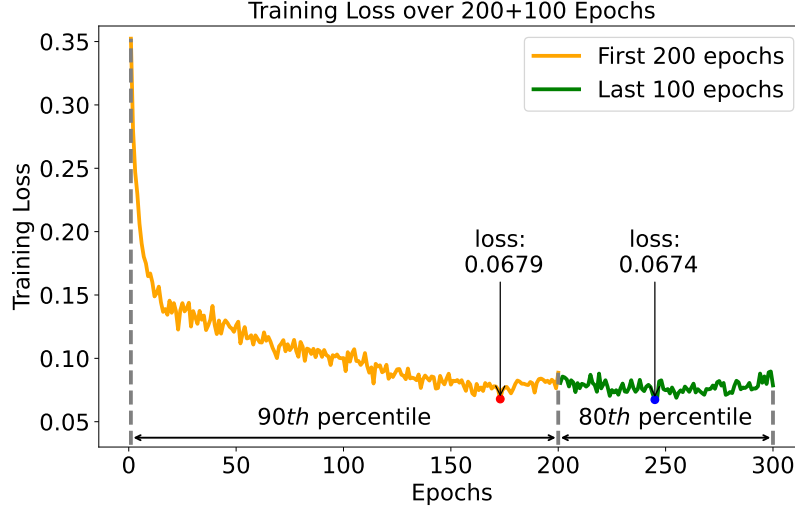


Figure 13: (best in color) Training loss of FoMo-0D ($D = 100$) with 0.5K unique + 7.5K transformed datasets/epoch for 200 epochs (in orange), followed with additional 100 epochs of training (in green). For the first 200 epochs we train with 90th percentile as the inlier/outlier threshold, which we reduce to 80th in the subsequent 100 epochs.

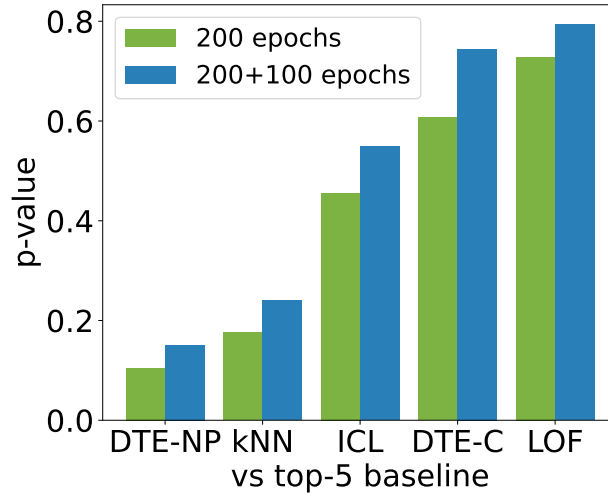


Figure 14: p -values increase with additional 100 epochs of pretraining, i.e. FoMo-0D w/ $D = 100$ performs better against top-5 baselines on datasets w/ $d \leq 100$.

We pretrain FoMo-0D on synthetic datasets from a simple data prior based on GMMs. The real-world benchmark datasets, on the other hand, may exhibit features with distributions different from Gaussians. To close the gap, we apply a quantile transform (denoted QT) on the benchmark datasets prior to feeding them to FoMo-0D for inference, which transforms the features to exhibit a more Gaussian-like probability distribution.

Figure 15 compares the performance of three FoMo-0D w/ $D = 100$ variants with and w/out QT against the top-5 baselines w.r.t. the p -values of the paired Wilcoxon signed rank test. FoMo-0D tends to perform better as suggested by larger p -values when QT is applied.

F.10 Effect of Model Size

To understand how the performance of FoMo-0D scales with model sizes, we vary the number of transformer layers $L = 1, 2, 3$, where the default is $L = 4$ (see in **Model architecture**). We present the p -values of

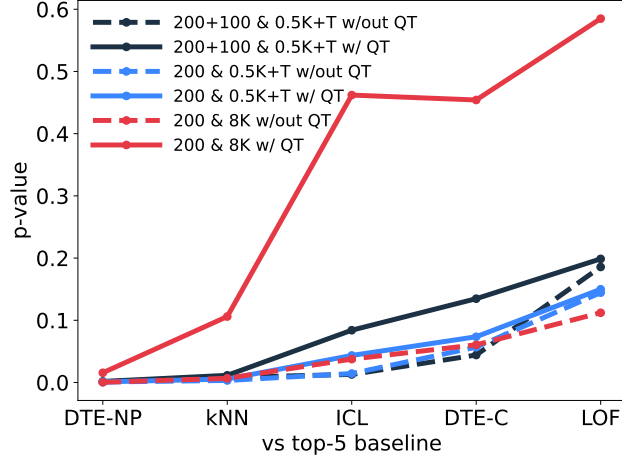


Figure 15: p -values increase, i.e. FoMo-0D performance improves, against top-5 baselines with quantile transform (QT) preceding inference, for 3 different settings of FoMo-0D w/ $D = 100$.

Table 8: p -values of the one-sided Wilcoxon signed rank test, comparing FoMo-0D (with $D = 100$) to **top 10 baselines** with default hyperparameters (HPs), and **top 4^{avg}** baselines⁶ with **avg.** performance over varying HPs (denoted w/ ^{avg}) over all (57) datasets. FoMo-0D improves (i.e. higher p -values) as number of transformer layers L increases. At $L = 1$, in-context learning ability appears to be quite limited. (setting: $D = 100$, $R = 500$, train/inference context size=5K, w/ quantile transform, $\alpha = 0.05$)

FoMo-0D	#parameters	DTE-NP	kNN	ICL	DTE-C	LOF	CBLOF	Feat.Bag.	SLAD	DDPM	OCSVM	DTE-NP ^{avg}	kNN ^{avg}	ICL ^{avg}	DTE-C ^{avg}
$L = 1$	1.34M	1.66×10^{-9}	4.30×10^{-9}	1.25×10^{-6}	1.34×10^{-7}	5.08×10^{-6}	5.44×10^{-8}	1.31×10^{-4}	1.07×10^{-6}	1.30×10^{-6}	1.46×10^{-7}	2.68×10^{-9}	1.53×10^{-8}	2.13×10^{-7}	0.395
$L = 2$	2.52M	0.006	0.036	0.157	0.259	0.442	0.557	0.703	0.431	0.759	0.805	0.021	0.134	0.333	1.000
$L = 3$	3.70M	0.016	0.098	0.372	0.579	0.572	0.871	0.808	0.652	0.921	0.961	0.085	0.335	0.652	1.000
$L = 4$	4.89M	0.016	0.106	0.462	0.454	0.585	0.750	0.823	0.759	0.901	0.895	0.112	0.315	0.670	1.000

different model sizes in Table 8, where a p -value > 0.05 means no statistical evidence to suggest performance difference between FoMo-0D and the compared baseline. We can observe that FoMo-0D is not comparable to the top-10 baselines with one layer, and shows improved performance (i.e., p -value increases and becomes larger than 0.05) as the number of layers increases, which suggests that scaling up the model size might help improve FoMo-0D’s performance. On the other hand, we can see that the detection performance didn’t increase a lot from $L = 3$ to 4, which suggests there exist diminishing returns in scaling the model size, and to further improve performance, one has to consider other methods (e.g., increase prior complexity, more pre-training epochs) besides scaling up only the model size.

G Generalization Analyses

G.1 Generalization to Out-of-Distribution Synthetic Datasets

We conduct analyses to understand FoMo’s ability to generalize on out-of-distribution synthetic GMM datasets. Besides the in-distribution setting for pre-training (i.e., $\mu \in [-5, 5]$, $\Sigma \in (0, 5]$, $m \leq 5$, $d \leq 100$), we consider the following out-of-distribution settings: (a) mean and covariance significantly out of range, with $\mu \in [-50, -5] \cup [5, 50]$, $\Sigma \in [5, 50]$, denoted as “ $|\mu|, |\Sigma| \in [5, 50]$ ”; (b) number of clusters significantly out of range, denoted as “ $m \in [5, 50]$ ”; (c) number of dimensions significantly out of range, denoted as “ $d \in [100, 500]$ ”; (d) binary outliers with values either 0 or 1 in one dimension from the sub-dimensions, denoted as “binary”; (e) “all”, which combines all the variants above. For each setting, we generate 1000 datasets with random seeds from 0 to 999, where on each dataset, we simulate 1000 test points with an outlier rate of 5% and evaluate FoMo-0D with a context length of 5000. We present the results with averaged performance over 1000 datasets for each setting in Table 9.

Table 9: Average metric score $\pm$ standard dev. over 1000 seeds for different out-of-distribution (OOD) synthetic GMMs. **FoMo-0D** remains robust against OOD test datasets as in (a)–(d), maintaining similar performance to in-distribution performance (top). Performance is affected more when datasets are OOD w.r.t. multiple factors combined as in (e).

Dataset	AUROC	AUCPR	F1
ID: in-distribution	98.55 $\pm$ 2.73	91.17 $\pm$ 13.07	86.74 $\pm$ 15.43
(a) OOD w.r.t. $ \mu , \Sigma \in [5, 50]$	94.79 $\pm$ 7.53	80.62 $\pm$ 21.19	76.32 $\pm$ 19.85
(b) OOD w.r.t. $m \in [5, 50]$	97.69 $\pm$ 3.59	86.72 $\pm$ 15.57	81.20 $\pm$ 16.23
(c) OOD w.r.t. $d \in [100, 500]$	96.22 $\pm$ 9.01	86.37 $\pm$ 23.27	83.23 $\pm$ 22.08
(d) OOD w.r.t. binary variable	100.00 $\pm$ 0.00	100.00 $\pm$ 0.06	99.97 $\pm$ 0.34
(e) OOD w.r.t. all combined	85.44 $\pm$ 16.96	64.17 $\pm$ 35.07	63.99 $\pm$ 33.53

We can observe different extents of performance degradation when applying out-of-distribution variations. Compared to other single variations, **FoMo-0D** seems to suffer more from inflating the mean and covariances, as due to the significant deviation in the parameters of the GMMs, inliers generated under such a setting are seemingly “outliers” w/o any reference points. Surprisingly, although **FoMo-0D** is only trained on continuous data, it can almost perfectly classify binary outliers hidden in one of the sub-dimensions, suggesting **FoMo-0D** could potentially generalize to discrete data at test time.

However, with all variations added, **FoMo-0D** becomes less capable compared to one single out-of-distribution variation, although there might exist some signals (e.g., binary labels) in favor of its decision-making process, for which training a powerful model with more comprehensive priors could possibly alleviate the issue.

G.2 Generalization to Out-of-GMM-Distribution Real-World Datasets

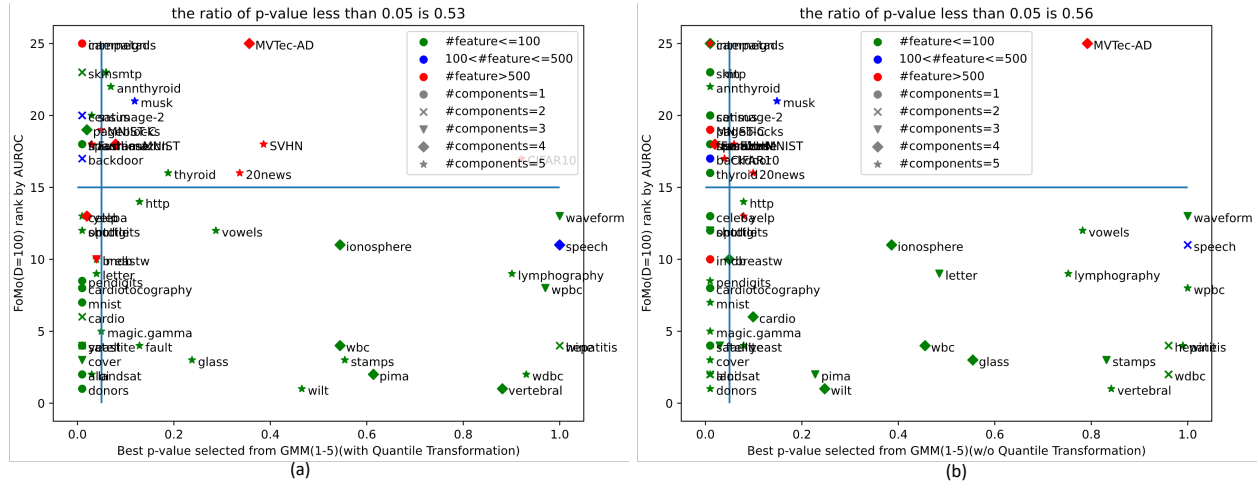


Figure 16: Goodness of fit test results for 57 datasets in ADBench. The x-axis is the p -value of a dataset, where a small p -value indicates GMMs are not a good fit for the dataset, and the y-axis is the rank of **FoMo-0D** ($D=100$) on that dataset (the smaller the better). We plot p -value = 0.05 (vertical) and rank = 15 (horizontal) as references. The left figure (a) is with quantile transformation while figure (b) is without quantile transformation. We use different colors to represent datasets at different dimensions, and use different markers to represent different numbers of clusters.

To understand how **FoMo-0D** generalizes from the pre-training GMM data priors to complex real-world datasets when performing zero-shot OD, we conduct the goodness of fit test Huber-Carol et al. (2012) for datasets in ADBench. We fit GMMs to each real-world dataset D_{real} with up to 5 components (as with our pre-training datasets), then sample D_{syn} from the best-parameter-fitted GMM and perform a two-sample

test¹⁰ on D_{real} and D_{syn} , with the null hypothesis that they come from the same distribution. A smaller p -value (≤ 0.05) of such a test provides evidence toward rejecting the null, which suggests GMM is not a good fit for the dataset (i.e., the pre-training data distribution is different from the test data).

We present the results in Figure 16, depicting the p -value (of the goodness-of-GMM-fit test) vs. FoMo-0D’s performance rank (the lower the better) among 30 baselines. We report the result both with quantile transformation in Figure 16(a) and without quantile transformation in 16(b). Since the two figures are highly similar, our next analysis will primarily focus on the figure with quantile transformation to align with our model’s implementation. We plot the vertical and horizontal lines as p -value = 0.05 and rank = 15. For p -value ≥ 0.05 and rank < 15 , we observe that performance is good on datasets with relatively large p -value where we cannot reject the null (i.e. GMM is a relatively good fit). This is where arguably FoMo-0D recalls its data prior distribution and generalizes to datasets similar to those seen during pretraining. We also see, for p -value < 0.05 and rank ≥ 15 , datasets with relatively poor performance where we can reject the null (i.e. GMM is not a good fit). These can be attributed to falling short in generalization to OOD datasets.

On the other hand, we observe many datasets concentrate on p -value < 0.05 and rank < 15 , where p -value is small (GMM not a good fit) yet the performance is competitive — those are the datasets on which FoMo-0D is likely to have achieved out-of-distribution generalization. It remains an open (theoretical) question to understand what (algorithm, if any) FoMo-0D might have learned that generalizes to out-of-distribution datasets. It is also an open (empirical) quest to explore whether a more complex data prior, beyond GMMs, could further push the performance up and by how much.

G.3 Generalization to Out of Distribution (OOD) Detection Tasks

We further evaluate FoMo-0D on more complex datasets (e.g., ImageNet-level). Specifically, we employ OpenOOD Zhang et al. (2023), an out-of-distribution (OOD) detection benchmark, where models are trained on labeled in-distribution datasets, with K known classes, and then evaluated on out-of-distribution datasets, aiming to detect $K + 1, K + 2, \dots$ novel classes. Although OOD detection is inherently different from OD, we can construct an OD dataset from OOD datasets, treating all K class samples as inliers and the $K + 1, K + 2, \dots$ OOD samples as outliers. For the in-distribution datasets, we choose ImageNet1K, which contains 1000 categories of images, and ImageNet200, a subset of ImageNet1K containing 200 categories. We further choose SSB-hard, NINCO, iNaturalist, Textures, and OpenImage-O as the out-of-distribution datasets, which gives us a total number of $2 \times 5 = 10$ datasets that are ImageNet-level complex.

Following Han et al. (2022), we create 10 new OD datasets from OpenOOD containing 10,000 samples with 5% outliers, and use the embedding from the last average pooling layer of ResNet18 He et al. (2016) as the feature (512) for each sample. Comparing FoMo-0D with the top-4 (on our original testbed) baselines in the order of: DTE-NP, k NN, ICL, DTE-C, we follow Livernoche et al. (2024) and report mean (standard dev.) over 5 runs (seed=0/1/2/3/4) on each dataset. We present the results with in-distribution datasets being ImageNet200 and ImageNet1K in Table 10 and 11, respectively.

Table 10: Average AUROC score $\pm$ standard dev. over five seeds for in-distribution dataset being **ImageNet200**. We use **blue** and **green** respectively to mark the **top-1** and the **top-2** method.

dataset	DTE-NP	kNN	ICL	DTE-C	FoMo-0D
ssb-hard	58.03 $\pm$ 0.00	58.14 $\pm$ 0.00	60.52 $\pm$ 0.25	60.74 $\pm$ 1.88	58.34 $\pm$ 1.55
ninco	53.28 $\pm$ 0.00	54.14 $\pm$ 0.00	59.56 $\pm$ 0.63	58.83 $\pm$ 1.54	55.16 $\pm$ 2.19
inaturalist	29.38 $\pm$ 0.00	29.51 $\pm$ 0.00	35.96 $\pm$ 1.10	41.77 $\pm$ 2.84	38.85 $\pm$ 3.29
textures	59.28 $\pm$ 0.00	59.91 $\pm$ 0.00	66.40 $\pm$ 0.69	70.33 $\pm$ 3.18	59.89 $\pm$ 2.07
openimageo	52.82 $\pm$ 0.00	53.79 $\pm$ 0.00	55.20 $\pm$ 0.69	59.09 $\pm$ 1.50	54.77 $\pm$ 1.19
average	50.56	51.10	55.53	58.15	53.40

We further report the p -value of the Wilcoxon signed rank test between the baselines and FoMo-0D on the 10 datasets from OpenOOD, as well as on the expanded benchmark combining those 10 with our original

¹⁰We use e-test from <https://www.rdocumentation.org/packages/energy/versions/1.7-11/topics/eqdist.etest>

Table 11: Average AUROC score $\pm$ standard dev. over five seeds for in-distribution dataset being **ImageNet1K**. We use **blue** and **green** respectively to mark the **top-1** and the **top-2** method.

dataset	DTE-NP	kNN	ICL	DTE-C	FoMo-0D
ssb-hard	55.63 $\pm$ 0.00	55.94 $\pm$ 0.00	58.79 $\pm$ 1.20	59.17 $\pm$ 1.82	56.73 $\pm$ 2.65
ninco	48.23 $\pm$ 0.00	49.10 $\pm$ 0.00	55.25 $\pm$ 0.87	57.60 $\pm$ 3.93	52.70 $\pm$ 2.70
inaturalist	30.24 $\pm$ 0.00	30.28 $\pm$ 0.00	35.03 $\pm$ 1.42	41.96 $\pm$ 3.13	38.94 $\pm$ 4.59
textures	54.38 $\pm$ 0.00	55.43 $\pm$ 0.00	61.30 $\pm$ 0.95	63.10 $\pm$ 3.72	55.18 $\pm$ 2.92
openimageo	54.31 $\pm$ 0.00	54.91 $\pm$ 0.00	54.02 $\pm$ 0.43	58.71 $\pm$ 2.08	56.95 $\pm$ 3.89
average	48.56	49.13	52.88	56.11	52.10

ADBench (10+57) in Table 12. In terms of metric values, FoMo-0D performs 2nd or 3rd best across OOD datasets. p -values show that it significantly outperforms DTE-NP and kNN ($p > 0.95$, such that the p -value < 0.05 for rejecting the null hypothesis and accepting the alternative hypothesis that the “baseline-minus-FoMo-0D” gap is smaller than zero) and is no different from ICL (2nd best after DTE-C). These results demonstrate that FoMo-0D generalizes beyond OD datasets and maintains strong zero-shot OD performance on complex, ImageNet-level OOD benchmarks.

Table 12: p -value of the Wilcoxon signed rank test (alternative: “greater”) between baselines and FoMo-0D on OpenOOD and combined benchmark on AUROC. A small p -value (≤ 0.05) means that there is statistical evidence for the alternative hypothesis such that baselines achieve higher metric performance than FoMo-0D.

method	DTE-NP	kNN	ICL	DTE-C
OpenOOD	1	0.9951	0.1875	0.0009
OpenOOD+ADBench	0.1271	0.3308	0.3153	0.1265

Interestingly, we observe that ICL and DTE-C outperform DTE-NP and kNN on the OpenOOD datasets, whereas on ADBench, DTE-NP and kNN are the top-2 methods outperforming ICL and DTE-C. We hypothesize this is because it is harder for non-parametric methods like DTE-NP and kNN to estimate meaningful decision boundaries in high dimensions (e.g., 512). In contrast, the performance of FoMo-0D is consistently competitive, where the p -values on the combined testbed (OpenOOD+ADBench) show that FoMo-0D is as competitive as all the top baselines across 67 diverse datasets, while maintaining zero-shot detection ability.

H Performance Profile Plots

To enable a comprehensive comparison of different methods, we plot rank (w.r.t. AUROC, lower is better) distribution in Figure 17, and adopt τ performance profile plots as described in Dolan & Moré (2002). These plots display the cumulative distribution of the τ metric—which quantifies suboptimality relative to the best-performing method. By computing sorted τ values along with their cumulative probabilities, we then use the area under each CDF curve as a global performance indicator, where a larger area signifies superior performance.

Figure 18, Figure 19, and Figure 20 illustrate performance profile plots of FoMo-0D and other baselines across all datasets. The results show that FoMo-0D (D=100) ranks at **top-5 (w.r.t. AUROC)**, **top-3 (w.r.t. AUPR)** and **top-1 (w.r.t. F1)**, respectively, outperforming many baselines.

Moreover, the performance of FoMo-0D (D=100) is even better (i.e., ranked within top-2) when tested on datasets with dimensions less than 100. As shown in Figure 21, Figure 22, and Figure 23, the area under the curve of FoMo-0D (D=100) ranks at **top-1 (w.r.t. AUROC)**, **top-2 (w.r.t. AUPR)** and **top-2 (w.r.t. F1)**, respectively, under this setting.

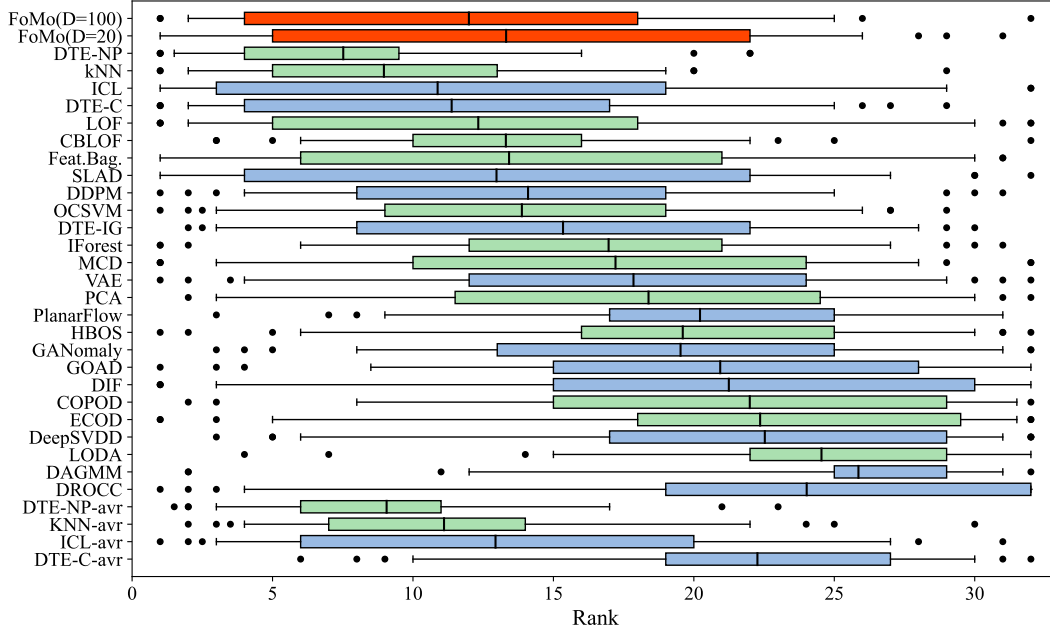


Figure 17: (best in color) Rank (w.r.t. AUROC, lower is better) distribution across all **57** real-world datasets shown via boxplots for (from top to bottom) FoMo-0D in **red**, all **26** baselines ordered by mean p -value⁶ (shallow and deep baselines in **green** and **blue**), and **top 4** baselines⁷ $_{avr}$ variants. The vertical line depicts the mean, the box shows the 25-75%, bars range 5-95%, and circles show the datasets at the tails.

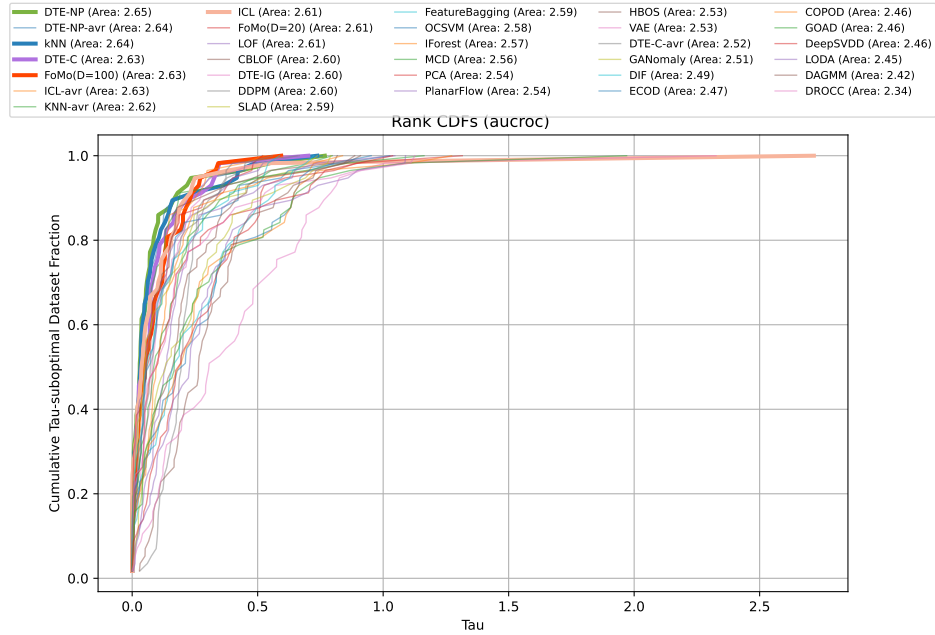


Figure 18: FoMo-0D ranks in **top-5** based on the performance profile plots of all detectors w.r.t. **AUROC** across **all datasets**. In the plot, x-axis represents the τ values—performance ratios that compare each method’s metric to the best performance achieved, while y-axis displays the cumulative fraction of test datasets for which a method’s performance is within the τ value. We use the area under each CDF curve as a global performance indicator, where a larger area signifies superior performance.

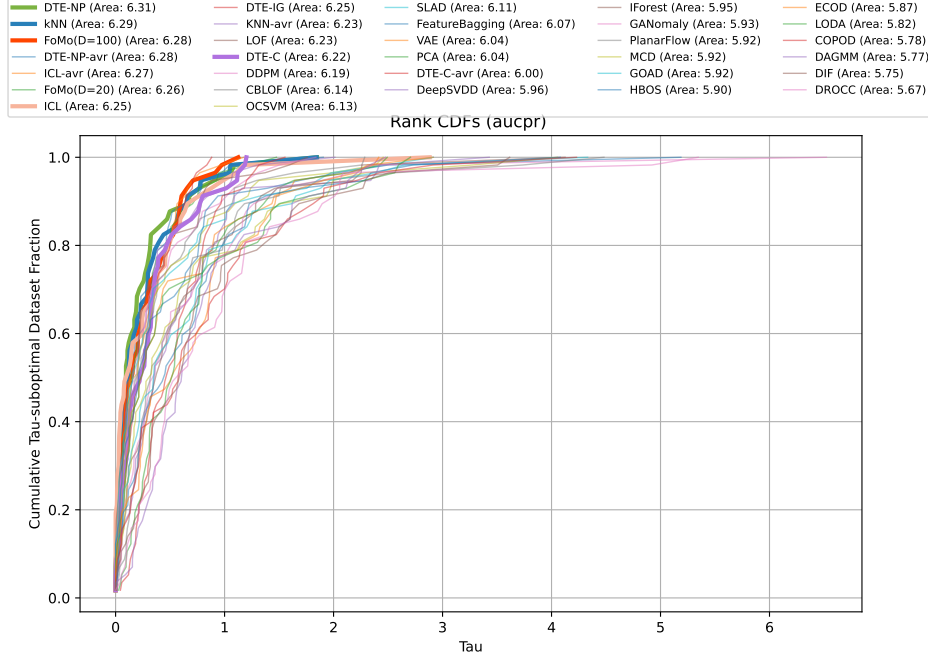


Figure 19: FoMo-0D ranks at **top-3** based on the performance profile plots of all detectors w.r.t. **AUPR** across **all datasets**. In the plot, x-axis represents the τ values—performance ratios that compare each method’s metric to the best performance achieved, while y-axis displays the cumulative fraction of test datasets for which a method’s performance is within the τ value. We use the area under each CDF curve as a global performance indicator, where a larger area signifies superior performance.

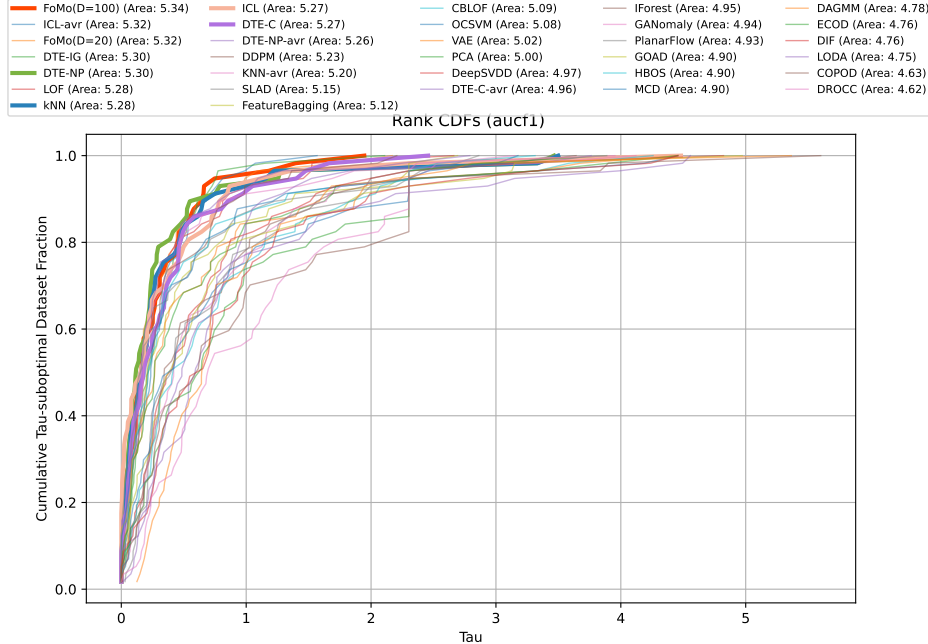


Figure 20: FoMo-0D ranks at **top-1** based on the performance profile plots of all detectors w.r.t. **F1** across all datasets. In the plot, x-axis represents the τ values—performance ratios that compare each method’s metric to the best performance achieved, while y-axis displays the cumulative fraction of test datasets for which a method’s performance is within the τ value. We use the area under each CDF curve as a global performance indicator, where a larger area signifies superior performance.

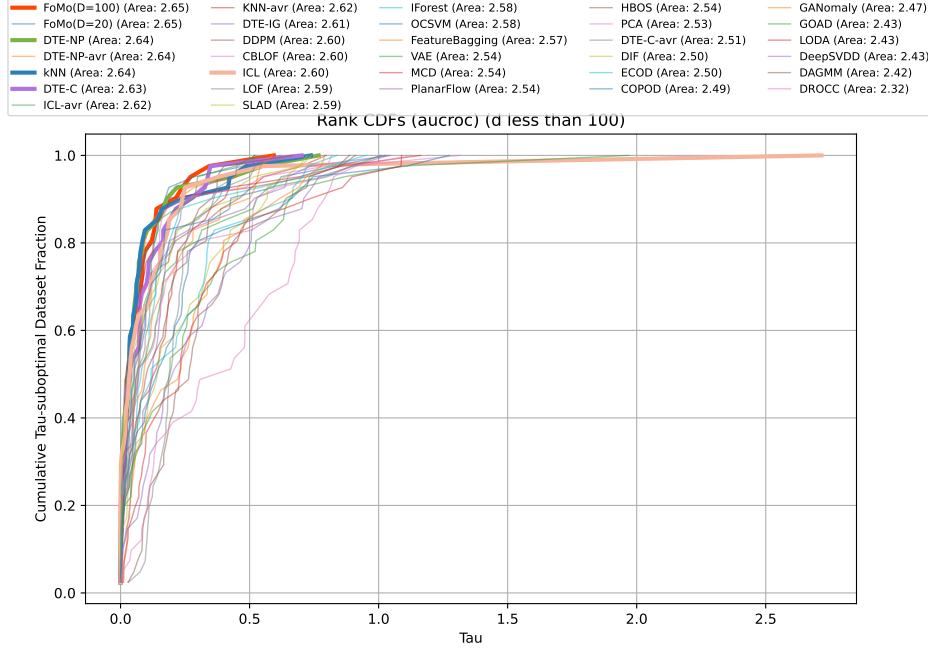


Figure 21: FoMo-0D ($D=100$) ranks at **top-1** based on the performance profile plots of all detectors w.r.t. **AU-ROC** in *datasets with dimensions less than $d \leq 100$* . In the plot, x-axis represents the τ values—performance ratios that compare each method’s metric to the best performance achieved, while y-axis displays the cumulative fraction of test datasets for which a method’s performance is within the τ value. We use the area under each CDF curve as a global performance indicator, where a larger area signifies superior performance.

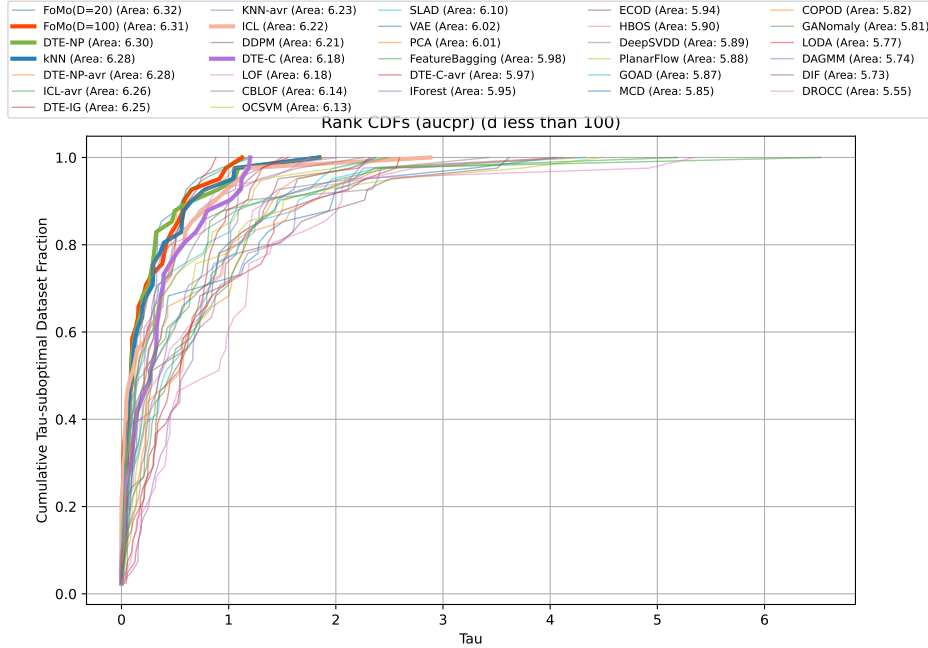


Figure 22: FoMo-0D ($D=100$) ranks at **top-2** based on the performance profile plots of all detectors w.r.t. **AUPR** in *datasets with dimensions less than $d \leq 100$* . In the plot, x-axis represents the τ values—performance ratios that compare each method’s metric to the best performance achieved, while y-axis displays the cumulative fraction of test datasets for which a method’s performance is within the τ value. We use the area under each CDF curve as a global performance indicator, where a larger area signifies superior performance.

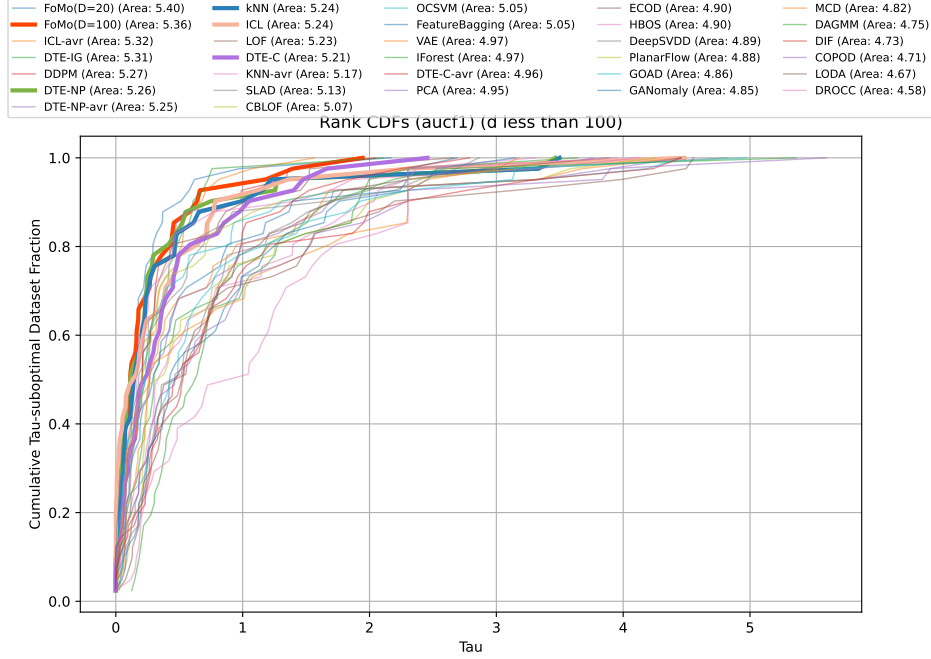


Figure 23: FoMo-0D ($D=100$) ranks at **top-2** based on the performance profile plots of all detectors w.r.t. **F1** in *datasets with dimensions less than $d \leq 100$* . In the plot, x-axis represents the τ values—performance ratios that compare each method’s metric to the best performance achieved, while y-axis displays the cumulative fraction of test datasets for which a method’s performance is within the τ value. We use the area under each CDF curve as a global performance indicator, where a larger area signifies superior performance.

I Full Results

Tables 14.1& 14.2, 15.1 & 15.2, and 16.1 & 16.2 respectively show the AUROC, AUPR, and F1 scores of the top-4 baselines, DTE-NP, kNN , ICL, and DTE-C as well as their corresponding avg model with the average performance across HPs, as listed in Table 4.

Tables 17.1&17.2, 18.1&18.2, and 19.1&19.2 respectively show the AUROC, AUPR, and F1 scores of all methods across all benchmark datasets. In all these tables, the last four rows show the avg_rank of methods across datasets, and p -values of the Wilcoxon signed rank test comparing FoMo-0D w/ $D = 100$ with other baselines. The preceding four rows are the same for FoMo-0D w/ $D = 20$, when ranking 31 models (26 baselines + 4 avg variants of top-4 baselines + FoMo-0D w/ $D = 20$).

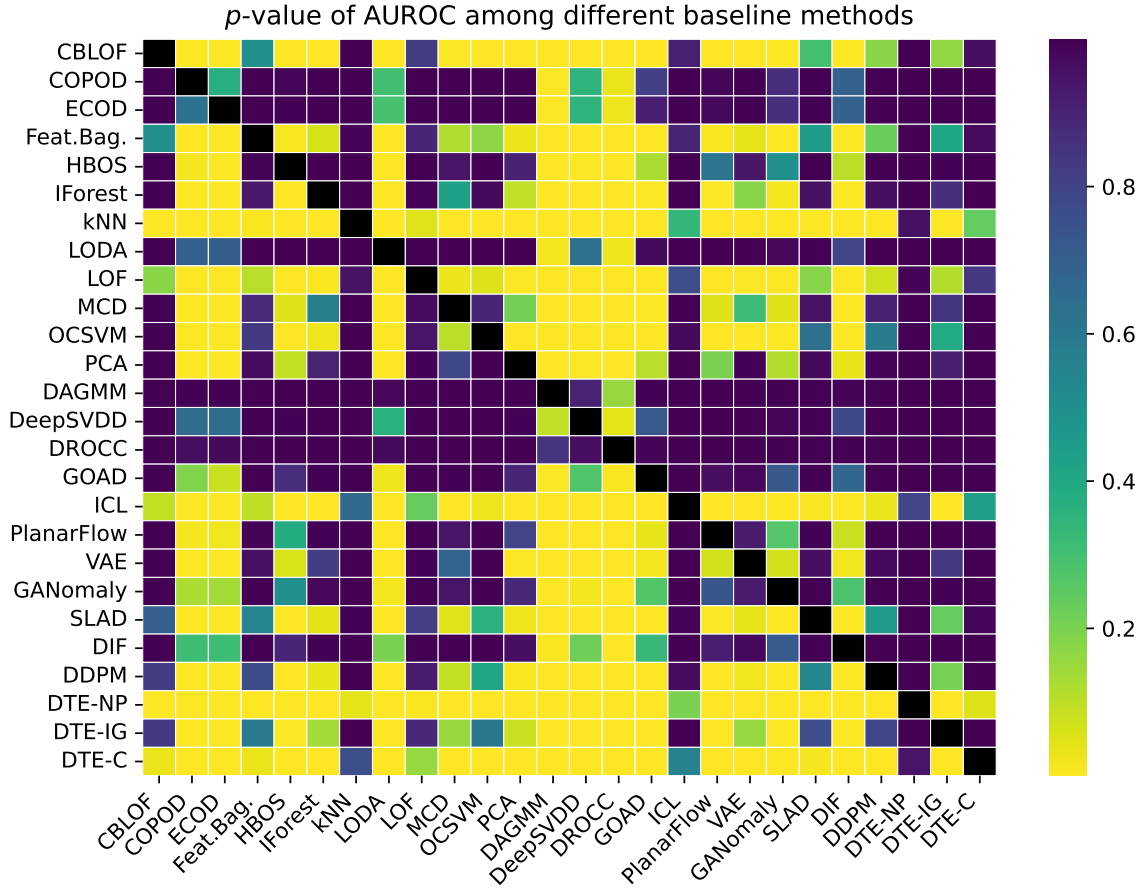


Figure 24: Pairwise *p*-values among baseline methods based on the Wilcoxon signed rank test w.r.t. AUROC performances across datasets.

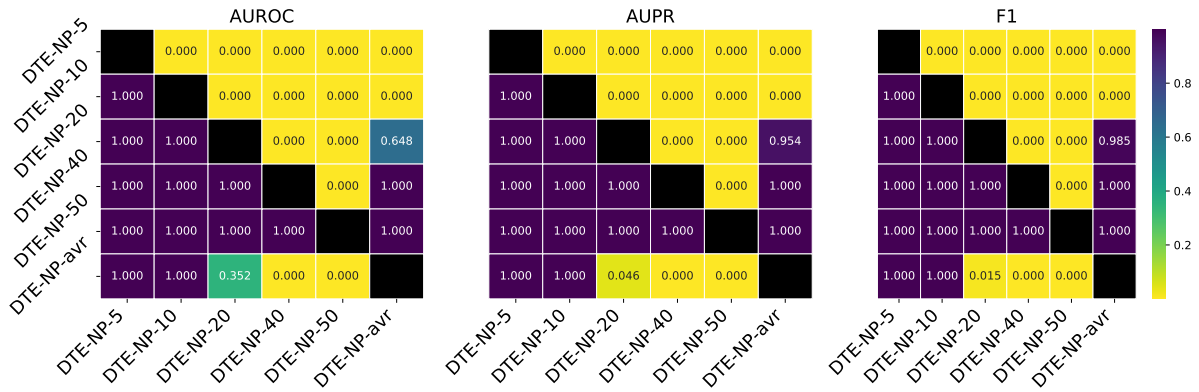


Figure 25: *p*-values w.r.t. AUROC/AUPR/F1 among different HP configurations of **DTE-NP** (i.e., $k \in \{5, 10, 20, 40, 50\}$), along with the ^{avg} model with the average performance across HPs.

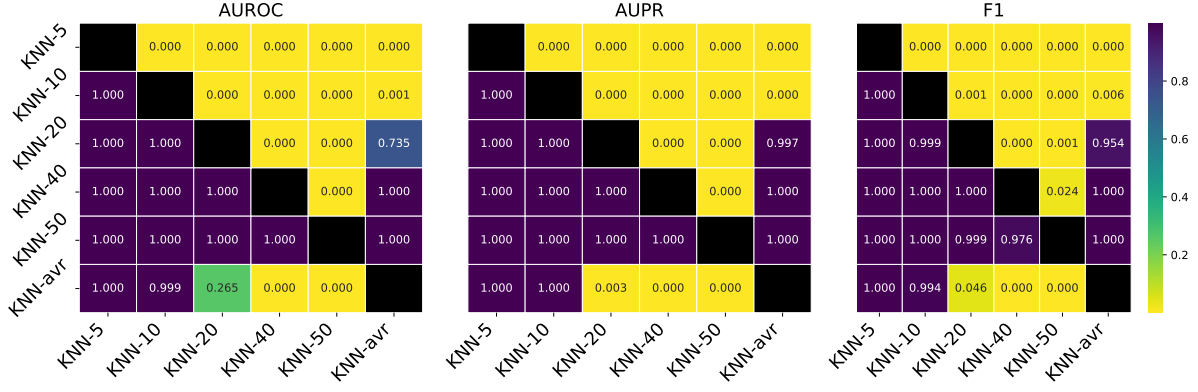


Figure 26: p -values w.r.t. AUROC/AUPR/F1 among different HP configurations of k NN (i.e., $k \in \{5, 10, 20, 40, 50\}$), along with the avr model with the average performance across HPs.

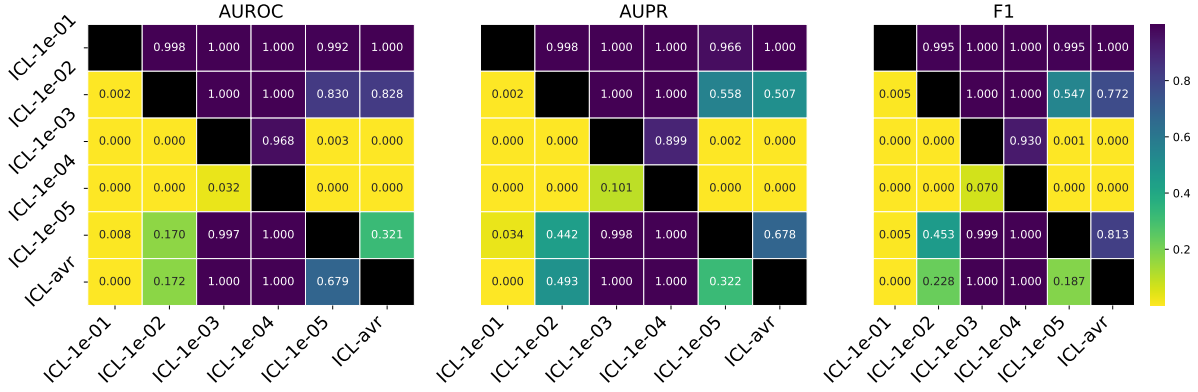


Figure 27: p -values w.r.t. AUROC/AUPR/F1 among different HP configurations of **ICL** (i.e., $\text{learning_rate} \in \{10^{-1}, 10^{-2}, 10^{-3}, 10^{-4}, 10^{-5}\}$), along with the avr model with the average performance across HPs.

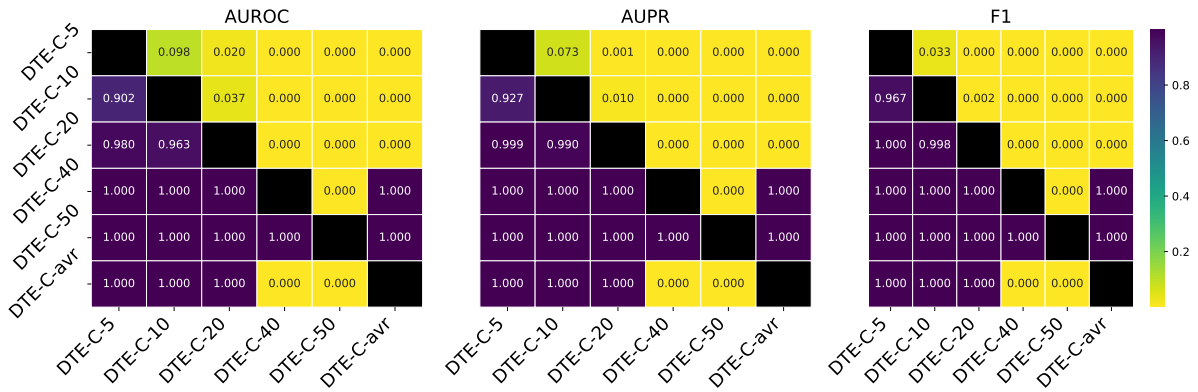


Figure 28: p -values w.r.t. AUROC/AUPR/F1 among different HP configurations of **DTE-C** (i.e., $k \in \{5, 10, 20, 40, 50\}$), along with the avr model with the average performance across HPs.

Table 13: Comparison of methods across datasets. (top row) Rank w.r.t. AUROC performance avg.'ed over 57 datasets is presented for FoMo-0D (with $D = 100$), **top-10 baselines** with default HPs, and **top-4**⁶ baselines with performance avg.'ed over varying HPs (denoted w/ ^{avg}); followed by p -values of the pairwise Wilcoxon signed rank test, comparing FoMo-0D to each baseline (from top to bottom) over All (57) datasets, those (24) w/ $d \leq 20$, (38) w/ $d \leq 50$, (42) w/ $d \leq 100$ and (46) datasets w/ $d \leq 500$ dimensions. FoMo-0D performs as well as (**i.e., statistically no different from**) **the 2nd best model** (k NN, w/ $p = 0.106$) across All datasets, while it is **comparable to** ($p > 0.05$) **or better than** ($p > 0.95$) **all baselines** over datasets w/ $d \leq 100$ (aligned w/ pretraining where $D = 100$) and $d \leq 500$ (generalizing beyond pretraining).

	FoMo-0D	DTE-NP	k NN	ICL	DTE-C	LOF	CBLOF	Feat.Bag.	SLAD	DDPM	OCSVM	DTE-NP ^{avg}	k NN ^{avg}	ICL ^{avg}	DTE-C ^{avg}
Rank(avg)	11.886	7.553	9.018	10.851	11.36	12.316	13.342	13.386	12.982	14.061	13.851	9.079	11.105	12.991	22.263
All	-	<u>0.016</u>	0.106	0.462	0.454	0.585	0.750	0.823	0.759	0.901	0.895	0.112	0.315	0.670	1.000
$d \leq 20$	-	0.428	0.665	0.987	0.727	0.911	0.940	0.987	0.868	0.758	0.968	0.781	0.868	0.990	1.000
$d \leq 50$	-	0.734	0.923	0.992	0.973	0.989	0.987	0.999	0.948	0.985	0.986	0.948	0.967	0.989	1.000
$d \leq 100$	-	0.415	0.700	0.949	0.953	0.970	0.971	0.996	0.876	0.980	0.978	0.752	0.860	0.958	1.000
$d \leq 200$	-	0.315	0.605	0.923	0.919	0.944	0.977	0.990	0.904	0.970	0.983	0.663	0.789	0.937	1.000
$d \leq 500$	-	0.220	0.569	0.827	0.894	0.960	0.968	0.994	0.910	0.960	0.979	0.607	0.756	0.846	1.000

Table 14.1: Average AUROC $\pm$ standard dev. over five seeds for the semi-supervised setting of DTE-NP, k NN with varying hyperparameter (HP) values; $k \in \{5, 10, 20, 40, 50\}$. Also reported is the avg model. We use **bold** and underline respectively to mark the **best** and the worst performance of each model to showcase the variability of performance across different HP settings.

dataset	DTE-NP-5	DTE-NP-10	DTE-NP-20	DTE-NP-40	DTE-NP-50	DTE-NP-avg	KNN-5	KNN-10	KNN-20	KNN-40	KNN-50	KNN-avg
amazon	50.69 $\pm$ 0.00	51.02 $\pm$ 0.00	51.26 $\pm$ 0.00	51.58 $\pm$ 0.00	51.69 $\pm$ 0.00	51.25 $\pm$ 0.00	51.04 $\pm$ 0.00	51.33 $\pm$ 0.00	51.63 $\pm$ 0.00	51.97 $\pm$ 0.00	52.08 $\pm$ 0.00	51.61 $\pm$ 0.00
amazon	60.76 $\pm$ 0.00	60.69 $\pm$ 0.00	60.53 $\pm$ 0.00	60.17 $\pm$ 0.00	60.22 $\pm$ 0.00	60.47 $\pm$ 0.00	60.58 $\pm$ 0.00	60.52 $\pm$ 0.00	60.23 $\pm$ 0.00	60.02 $\pm$ 0.00	60.25 $\pm$ 0.00	60.25 $\pm$ 0.00
amthlyroid	93.01 $\pm$ 0.00	92.89 $\pm$ 0.00	92.66 $\pm$ 0.00	92.30 $\pm$ 0.00	92.28 $\pm$ 0.00	92.65 $\pm$ 0.00	92.81 $\pm$ 0.00	92.60 $\pm$ 0.00	92.44 $\pm$ 0.00	91.98 $\pm$ 0.00	91.79 $\pm$ 0.00	92.30 $\pm$ 0.00
backdoor	94.48 $\pm$ 0.42	93.72 $\pm$ 0.46	92.67 $\pm$ 0.46	91.20 $\pm$ 0.45	92.08 $\pm$ 0.45	92.55 $\pm$ 0.44	92.81 $\pm$ 0.46	92.58 $\pm$ 0.46	91.14 $\pm$ 0.46	89.18 $\pm$ 0.47	88.39 $\pm$ 0.51	91.00 $\pm$ 0.47
breastw	99.10 $\pm$ 0.28	98.91 $\pm$ 0.35	98.59 $\pm$ 0.34	98.40 $\pm$ 0.35	98.36 $\pm$ 0.28	98.67 $\pm$ 0.28	99.09 $\pm$ 0.24	99.11 $\pm$ 0.27	99.16 $\pm$ 0.22	99.21 $\pm$ 0.18	99.21 $\pm$ 0.17	99.16 $\pm$ 0.21
campaign	78.34 $\pm$ 0.00	78.71 $\pm$ 0.00	78.91 $\pm$ 0.00	78.93 $\pm$ 0.00	78.90 $\pm$ 0.00	78.76 $\pm$ 0.00	78.48 $\pm$ 0.00	78.74 $\pm$ 0.00	78.84 $\pm$ 0.00	78.65 $\pm$ 0.00	78.66 $\pm$ 0.00	78.66 $\pm$ 0.00
cardio	91.33 $\pm$ 0.00	92.03 $\pm$ 0.00	92.46 $\pm$ 0.00	93.06 $\pm$ 0.00	93.28 $\pm$ 0.00	92.47 $\pm$ 0.00	92.00 $\pm$ 0.00	92.44 $\pm$ 0.00	92.99 $\pm$ 0.00	97.85 $\pm$ 0.00	94.08 $\pm$ 0.00	93.07 $\pm$ 0.00
cardiolography	60.40 $\pm$ 0.00	61.63 $\pm$ 0.00	63.14 $\pm$ 0.00	65.05 $\pm$ 0.00	65.81 $\pm$ 0.00	63.21 $\pm$ 0.00	62.11 $\pm$ 0.00	63.39 $\pm$ 0.00	65.12 $\pm$ 0.00	67.85 $\pm$ 0.00	68.91 $\pm$ 0.00	65.48 $\pm$ 0.00
celeba	72.18 $\pm$ 0.34	72.34 $\pm$ 0.17	72.28 $\pm$ 0.10	71.93 $\pm$ 0.17	71.89 $\pm$ 0.17	72.42 $\pm$ 0.28	72.91 $\pm$ 0.29	72.36 $\pm$ 0.12	71.94 $\pm$ 0.16	71.37 $\pm$ 0.21	71.28 $\pm$ 0.16	71.84 $\pm$ 0.15
census	97.90 $\pm$ 0.17	97.72 $\pm$ 0.14	97.40 $\pm$ 0.18	96.99 $\pm$ 0.23	96.84 $\pm$ 0.24	97.37 $\pm$ 0.19	97.51 $\pm$ 0.15	97.19 $\pm$ 0.15	96.75 $\pm$ 0.22	96.21 $\pm$ 0.28	96.00 $\pm$ 0.16	96.73 $\pm$ 0.22
cover	99.72 $\pm$ 0.03	99.61 $\pm$ 0.03	99.43 $\pm$ 0.06	99.14 $\pm$ 0.09	99.02 $\pm$ 0.10	99.38 $\pm$ 0.06	99.51 $\pm$ 0.06	99.24 $\pm$ 0.08	98.20 $\pm$ 0.13	97.30 $\pm$ 0.14	97.90 $\pm$ 0.11	98.74 $\pm$ 0.09
donors	58.31 $\pm$ 0.00	58.37 $\pm$ 0.03	58.70 $\pm$ 0.03	60.00 $\pm$ 0.00	60.43 $\pm$ 0.00	59.17 $\pm$ 0.02	58.73 $\pm$ 0.00	58.76 $\pm$ 0.00	60.12 $\pm$ 0.00	61.71 $\pm$ 0.00	61.79 $\pm$ 0.00	60.22 $\pm$ 0.00
fraud	95.70 $\pm$ 0.90	95.67 $\pm$ 0.93	95.64 $\pm$ 0.93	95.60 $\pm$ 0.92	95.60 $\pm$ 0.92	95.60 $\pm$ 0.92	95.59 $\pm$ 0.97	95.55 $\pm$ 0.99	95.55 $\pm$ 0.89	95.54 $\pm$ 0.92	95.62 $\pm$ 0.88	95.57 $\pm$ 0.93
glass	96.08 $\pm$ 0.39	93.04 $\pm$ 1.06	89.82 $\pm$ 1.12	87.89 $\pm$ 1.10	81.97 $\pm$ 1.40	90.83 $\pm$ 0.91	92.18 $\pm$ 0.94	88.67 $\pm$ 0.98	87.24 $\pm$ 1.18	84.93 $\pm$ 2.92	83.55 $\pm$ 2.61	87.39 $\pm$ 1.59
hepatitis	99.84 $\pm$ 0.20	99.27 $\pm$ 0.51	96.89 $\pm$ 0.96	93.15 $\pm$ 1.69	91.36 $\pm$ 1.76	96.22 $\pm$ 0.88	96.77 $\pm$ 1.47	86.88 $\pm$ 2.21	85.20 $\pm$ 2.34	85.46 $\pm$ 1.92	84.88 $\pm$ 2.09	87.90 $\pm$ 1.75
lftp	99.99 $\pm$ 0.00	99.98 $\pm$ 0.01	99.95 $\pm$ 0.00	99.93 $\pm$ 0.01	99.95 $\pm$ 0.01	99.95 $\pm$ 0.01	100.00 $\pm$ 0.00	99.94 $\pm$ 0.02	99.95 $\pm$ 0.01	99.95 $\pm$ 0.01	99.95 $\pm$ 0.01	99.96 $\pm$ 0.01
indb	50.48 $\pm$ 0.00	50.38 $\pm$ 0.00	50.32 $\pm$ 0.00	50.28 $\pm$ 0.00	50.27 $\pm$ 0.00	50.27 $\pm$ 0.00	50.08 $\pm$ 0.00	50.00 $\pm$ 0.00	50.29 $\pm$ 0.00	50.23 $\pm$ 0.00	50.23 $\pm$ 0.00	50.18 $\pm$ 0.00
intercards	70.96 $\pm$ 0.00	68.65 $\pm$ 0.00	66.86 $\pm$ 0.00	65.97 $\pm$ 0.00	65.82 $\pm$ 0.00	67.65 $\pm$ 0.00	68.08 $\pm$ 0.00	65.48 $\pm$ 0.00	65.04 $\pm$ 0.00	65.04 $\pm$ 0.00	65.04 $\pm$ 0.00	65.73 $\pm$ 0.00
ionosphere	98.48 $\pm$ 0.60	98.13 $\pm$ 0.71	97.84 $\pm$ 0.64	96.83 $\pm$ 0.71	96.21 $\pm$ 0.70	97.50 $\pm$ 0.63	97.32 $\pm$ 0.85	97.02 $\pm$ 0.81	96.03 $\pm$ 0.76	92.80 $\pm$ 1.64	91.53 $\pm$ 1.65	95.12 $\pm$ 0.92
landsat	68.99 $\pm$ 0.00	68.02 $\pm$ 0.00	66.46 $\pm$ 0.00	64.73 $\pm$ 0.00	64.16 $\pm$ 0.00	66.47 $\pm$ 0.00	68.25 $\pm$ 0.00	66.48 $\pm$ 0.00	64.96 $\pm$ 0.00	62.49 $\pm$ 0.00	61.03 $\pm$ 0.00	64.70 $\pm$ 0.00
letter	36.12 $\pm$ 0.00	35.66 $\pm$ 0.00	34.78 $\pm$ 0.00	33.72 $\pm$ 0.00	33.09 $\pm$ 0.00	34.74 $\pm$ 0.00	35.43 $\pm$ 0.00	34.54 $\pm$ 0.00	33.17 $\pm$ 0.00	32.11 $\pm$ 0.00	31.69 $\pm$ 0.00	33.39 $\pm$ 0.00
lymphography	99.88 $\pm$ 0.25	99.70 $\pm$ 0.32	99.79 $\pm$ 0.32	99.76 $\pm$ 0.31	99.76 $\pm$ 0.31	99.80 $\pm$ 0.30	99.87 $\pm$ 0.10	99.85 $\pm$ 0.05	99.85 $\pm$ 0.05	99.88 $\pm$ 0.08	99.88 $\pm$ 0.08	99.87 $\pm$ 0.06
magic gamma	83.91 $\pm$ 0.00	83.40 $\pm$ 0.00	82.87 $\pm$ 0.00	82.05 $\pm$ 0.00	81.72 $\pm$ 0.00	82.81 $\pm$ 0.00	83.27 $\pm$ 0.00	82.81 $\pm$ 0.00	81.85 $\pm$ 0.00	80.76 $\pm$ 0.00	80.30 $\pm$ 0.00	81.76 $\pm$ 0.00
manimagraphy	87.65 $\pm$ 0.00	87.73 $\pm$ 0.00	87.68 $\pm$ 0.00	87.62 $\pm$ 0.00	87.29 $\pm$ 0.00	87.55 $\pm$ 0.00	87.53 $\pm$ 0.00	87.75 $\pm$ 0.00	87.38 $\pm$ 0.00	86.07 $\pm$ 0.00	86.76 $\pm$ 0.00	87.29 $\pm$ 0.00
mnist	94.22 $\pm$ 0.00	93.93 $\pm$ 0.00	93.57 $\pm$ 0.00	93.20 $\pm$ 0.00	93.05 $\pm$ 0.00	93.05 $\pm$ 0.00	93.85 $\pm$ 0.00	93.85 $\pm$ 0.00	93.05 $\pm$ 0.00	92.55 $\pm$ 0.00	92.26 $\pm$ 0.00	93.04 $\pm$ 0.00
mnist	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00
mnist	95.00 $\pm$ 0.00	93.97 $\pm$ 0.00	93.52 $\pm$ 0.00	90.87 $\pm$ 0.00	90.28 $\pm$ 0.00	92.33 $\pm$ 0.00	93.72 $\pm$ 0.00	92.69 $\pm$ 0.00	90.26 $\pm$ 0.00	87.66 $\pm$ 0.00	86.07 $\pm$ 0.00	90.07 $\pm$ 0.00
mnist	93.01 $\pm$ 0.00	93.40 $\pm$ 0.00	93.56 $\pm$ 0.00	93.37 $\pm$ 0.00	92.83 $\pm$ 0.00	93.42 $\pm$ 0.00	93.65 $\pm$ 0.00	93.86 $\pm$ 0.00	90.86 $\pm$ 0.00	89.30 $\pm$ 0.00	88.13 $\pm$ 0.00	89.57 $\pm$ 0.00
mnist	93.01 $\pm$ 0.00	93.40 $\pm$ 0.00	93.56 $\pm$ 0.00	93.37 $\pm$ 0.00	92.83 $\pm$ 0.00	93.42 $\pm$ 0.00	93.65 $\pm$ 0.00	93.86 $\pm$ 0.00	90.86 $\pm$ 0.00	89.30 $\pm$ 0.00	88.13 $\pm$ 0.00	89.57 $\pm$ 0.00
mnist	93.01 $\pm$ 0.00	93.40 $\pm$ 0.00	93.56 $\pm$ 0.00	93.37 $\pm$ 0.00	92.83 $\pm$ 0.00	93.42 $\pm$ 0.00	93.65 $\pm$ 0.00	93.86 $\pm$ 0.00	90.86 $\pm$ 0.00	89.30 $\pm$ 0.00	88.13 $\pm$ 0.00	89.57 $\pm$ 0.00
mnist	93.01 $\pm$ 0.00	93.40 $\pm$ 0.00	93.56 $\pm$ 0.00	93.37 $\pm$ 0.00	92.83 $\pm$ 0.00	93.42 $\pm$ 0.00	93.65 $\pm$ 0.00	93.86 $\pm$ 0.00	90.86 $\pm$ 0.00	89.30 $\pm$ 0.00	88.13 $\pm$ 0.00	89.57 $\pm$ 0.00
mnist	93.01 $\pm$ 0.00	93.40 $\pm$ 0.00	93.56 $\pm$ 0.00	93.37 $\pm$ 0.00	92.83 $\pm$ 0.00	93.42 $\pm$ 0.00	93.65 $\pm$ 0.00	93.86 $\pm$ 0.00	90.86 $\pm$ 0.00	89.30 $\pm$ 0.00	88.13 $\pm$ 0.00	89.57 $\pm$ 0.00
mnist	93.01 $\pm$ 0.00	93.40 $\pm$ 0.00	93.56 $\pm$ 0.00	93.37 $\pm$ 0.00	92.83 $\pm$ 0.00	93.42 $\pm$ 0.00	93.65 $\pm$ 0.00	93.86 $\pm$ 0.00	90.86 $\pm$ 0.00	89.30 $\pm$ 0.00	88.13 $\pm$ 0.00	89.57 $\pm$ 0.00
mnist	93.01 $\pm$ 0.00	93.40 $\pm$ 0.00	93.56 $\pm$ 0.00	93.37 $\pm$ 0.00	92.83 $\pm$ 0.00	93.42 $\pm$ 0.00	93.65 $\pm$ 0.00	93.86 $\pm$ 0.00	90.86 $\pm$ 0.00	89.30 $\pm$ 0.00	88.13 $\pm$ 0.00	89.57 $\pm$ 0.00
mnist	93.01 $\pm$ 0.00	93.40 $\pm$ 0.00	93.56 $\pm$ 0.00	93.37 $\pm$ 0.00	92.83 $\pm$ 0.00	93.42 $\pm$ 0.00	93.65 $\pm$ 0.00	93.86 $\pm$ 0.00	90.86 $\pm$ 0.00	89.30 $\pm$ 0.00	88.13 $\pm$ 0.00	89.57 $\pm$ 0.00
mnist	93.01 $\pm$ 0.00	93.40 $\pm$ 0.00	93.56 $\pm$ 0.00	93.37 $\pm$ 0.00	92.83 $\pm$ 0.00	93.42 $\pm$ 0.00	93.65 $\pm$ 0.00	93.86 $\pm$ 0.00	90.86 $\pm$ 0.00	89.30 $\pm$ 0.00	88.13 $\pm$ 0.00	89.57 $\pm$ 0.00
mnist	93.01 $\pm$ 0.00	93.40 $\pm$ 0.00	93.56 $\pm$ 0.00	93.37 $\pm$ 0.00	92.83 $\pm$ 0.00	93.42 $\pm$ 0.00	93.65 $\pm$ 0.00	93.86 $\pm$ 0.00	90.86 $\pm$ 0.00	89.30 $\pm$ 0.00	88.13 $\pm$ 0.00	89.57 $\pm$ 0.00
mnist	93.01 $\pm$ 0.00	93.40 $\pm$ 0.00	93.56 $\pm$ 0.00	93.37 $\pm$ 0.00	92.83 $\pm$ 0.00	93.42 $\pm$ 0.00	93.65 $\pm$ 0.00	93.86 $\pm$ 0.00	90.86 $\pm$ 0.00	89.30 $\pm$ 0.00	88.13 $\pm$ 0.00	89.57 $\pm$ 0.00
mnist	93.01 $\pm$ 0.00	93.40 $\pm$ 0.00	93.56 $\pm$ 0.00	93.37 $\pm$ 0.00	92.83 $\pm$ 0.00	93.42 $\pm$ 0.00	93.65 $\pm$ 0.00	93.86 $\pm$ 0.00	90.86 $\pm$ 0.00	89.30 $\pm$ 0.00	88.13 $\pm$ 0.00	89.57 $\pm$ 0.00
mnist	93.01 $\pm$ 0.00	93.40 $\pm$ 0.00	93.56 $\pm$ 0.00	93.37 $\pm$ 0.00	92.83 $\pm$ 0.00	93.42 $\pm$ 0.00	93.65 $\pm$ 0.00	93.86 $\pm$ 0.00	90.86 $\pm$ 0.00	89.30 $\pm$ 0.00	88.13 $\pm$ 0.00	89.57 $\pm$ 0.00
mnist	93.01 $\pm$ 0.00	93.40 $\pm$ 0.00	93.56 $\pm$ 0.00	93.37 $\pm$ 0.00	92.83 $\pm$ 0.00	93.42 $\pm$ 0.00	93.65 $\pm$ 0.00	93.86 $\pm$ 0.00	90.86 $\pm$ 0.00	89.30 $\pm$ 0.00	88.13 $\pm$ 0.00	89.57 $\pm$ 0.00
mnist	93.01 $\pm$ 0.00	93.40 $\pm$ 0.00	93.56 $\pm$ 0.00	93.37 $\pm$ 0.00	92.83 $\pm$ 0.00	93.42 $\pm$ 0.00	93.65 $\pm$ 0.00	93.86 $\pm$ 0.00	90.86 $\pm$ 0.00	89.30 $\pm$ 0.00	88.13 $\pm$ 0.00	89.57 $\pm$ 0.00
mnist	93.01 $\pm$ 0.00	93.40 $\pm$ 0.00	93.56 $\pm$ 0.00	93.37 $\pm$ 0.00	92.83 $\pm$ 0.00	93.42 $\pm$ 0.00	93.65 $\pm$ 0.00	93.86 $\pm$ 0.00	90.86 $\pm$ 0.00	89.30 $\pm$ 0.00	88.13 $\pm$ 0.00	89.57 $\pm$ 0.00
mnist	93.01 $\pm$ 0.00	93.40 $\pm$ 0.00	93.56 $\pm$ 0.00	93.37 $\pm$ 0.00	92.83 $\pm$ 0.00	93.42 $\pm$ 0.00	93.65 $\pm$ 0.00	93.86 $\pm$ 0.00	90.86 $\pm$ 0.00	89.30 $\pm$ 0.00</		

Table 14.2: Average AUROC $\pm$ standard dev. over five seeds for the semi-supervised setting of ICL and DTE-C baselines with varying hyperparameter (HP) values; For ICL, the learning rate $\in \{0.1, 0.02, 0.001, 0.0001, 1e-05\}$, for DTE-C, $k \in \{5, 10, 20, 40, 50\}$. Also reported is the avg model. We use **bold** and underline performance respectively to mark the **best** and the **worst** performance of each model to showcase the variability of performance across different HP settings.

dataset	ICL-0.1	ICL-0.01	ICL-0.001	ICL-0.0001	ICL-1e-05	ICL-avg	DTE-C-5	DTE-C-10	DTE-C-20	DTE-C-40	DTE-C-50	DTE-C-avg
aida	47.74 $\pm$ 0.28	47.12 $\pm$ 0.51	46.81 $\pm$ 0.47	48.42 $\pm$ 0.24	48.06 $\pm$ 0.23	47.63 $\pm$ 0.15	50.20 $\pm$ 0.21	50.84 $\pm$ 0.10	50.16 $\pm$ 0.34	50.26 $\pm$ 0.33	50.00 $\pm$ 0.00	50.29 $\pm$ 0.09
amazon	53.07 $\pm$ 0.19	53.44 $\pm$ 0.19	52.72 $\pm$ 0.68	53.31 $\pm$ 0.21	53.18 $\pm$ 0.10	53.15 $\pm$ 0.19	56.20 $\pm$ 0.12	55.07 $\pm$ 0.20	55.07 $\pm$ 0.20	55.00 $\pm$ 0.00	55.00 $\pm$ 0.00	53.48 $\pm$ 1.33
amulhyroid	84.02 $\pm$ 0.46	72.68 $\pm$ 3.79	72.29 $\pm$ 2.69	88.84 $\pm$ 2.35	88.52 $\pm$ 1.40	84.27 $\pm$ 0.10	97.47 $\pm$ 0.10	97.65 $\pm$ 0.11	97.73 $\pm$ 0.15	97.40 $\pm$ 0.22	50.00 $\pm$ 0.00	88.05 $\pm$ 0.05
backdoor	93.03 $\pm$ 0.66	<u>92.01</u> $\pm$ 0.70	93.30 $\pm$ 0.44	93.93 $\pm$ 0.58	93.32 $\pm$ 0.71	93.30 $\pm$ 0.48	98.65 $\pm$ 1.08	92.00 $\pm$ 1.04	92.64 $\pm$ 0.84	50.00 $\pm$ 0.00	50.00 $\pm$ 0.00	74.67 $\pm$ 0.36
breastw	98.96 $\pm$ 0.47	98.87 $\pm$ 0.20	99.19 $\pm$ 0.34	99.11 $\pm$ 0.18	97.61 $\pm$ 0.65	98.75 $\pm$ 1.34	93.45 $\pm$ 1.34	94.31 $\pm$ 1.25	97.41 $\pm$ 0.91	98.85 $\pm$ 0.45	99.26 $\pm$ 0.07	96.52 $\pm$ 0.53
campaign	76.07 $\pm$ 0.87	<u>74.61</u> $\pm$ 1.97	78.88 $\pm$ 0.42	79.79 $\pm$ 0.91	82.30 $\pm$ 0.38	78.33 $\pm$ 0.52	79.18 $\pm$ 1.09	78.24 $\pm$ 2.09	78.49 $\pm$ 1.13	50.00 $\pm$ 0.00	50.00 $\pm$ 0.00	67.18 $\pm$ 0.74
cardio	73.99 $\pm$ 7.79	84.98 $\pm$ 4.75	82.30 $\pm$ 3.36	78.68 $\pm$ 2.18	68.59 $\pm$ 0.64	77.71 $\pm$ 1.76	88.07 $\pm$ 0.51	60.09 $\pm$ 2.19	87.66 $\pm$ 0.63	50.00 $\pm$ 0.00	50.00 $\pm$ 0.00	72.66 $\pm$ 0.21
cardiologocography	48.99 $\pm$ 2.02	54.18 $\pm$ 2.77	53.24 $\pm$ 3.25	50.67 $\pm$ 2.55	47.20 $\pm$ 1.33	50.85 $\pm$ 1.24	60.36 $\pm$ 1.54	60.36 $\pm$ 1.54	50.00 $\pm$ 0.00	50.00 $\pm$ 0.00	50.00 $\pm$ 0.00	55.90 $\pm$ 0.23
celeba	79.38 $\pm$ 2.17	70.15 $\pm$ 2.47	76.13 $\pm$ 1.88	78.39 $\pm$ 1.78	79.43 $\pm$ 1.44	72.25 $\pm$ 1.63	82.95 $\pm$ 1.26	81.59 $\pm$ 0.79	80.16 $\pm$ 1.26	50.00 $\pm$ 0.00	50.00 $\pm$ 0.00	68.94 $\pm$ 0.44
census	70.41 $\pm$ 2.07	67.52 $\pm$ 8.79	73.98 $\pm$ 0.78	75.03 $\pm$ 0.46	74.37 $\pm$ 0.53	72.25 $\pm$ 1.63	82.95 $\pm$ 1.26	81.59 $\pm$ 0.79	80.16 $\pm$ 1.26	50.00 $\pm$ 0.00	50.00 $\pm$ 0.00	61.41 $\pm$ 0.71
cover	93.59 $\pm$ 3.09	92.32 $\pm$ 0.14	91.92 $\pm$ 4.58	94.97 $\pm$ 2.91	94.27 $\pm$ 3.40	91.42 $\pm$ 1.74	97.57 $\pm$ 0.86	97.81 $\pm$ 0.75	96.61 $\pm$ 0.63	50.00 $\pm$ 0.00	50.00 $\pm$ 0.00	78.40 $\pm$ 0.14
donors	86.20 $\pm$ 10.29	97.76 $\pm$ 1.57	98.64 $\pm$ 0.55	99.37 $\pm$ 0.30	99.51 $\pm$ 0.14	96.30 $\pm$ 1.95	98.68 $\pm$ 0.15	97.56 $\pm$ 0.53	95.64 $\pm$ 0.27	50.00 $\pm$ 0.00	50.00 $\pm$ 0.00	80.17 $\pm$ 3.55
fault	63.37 $\pm$ 3.29	63.04 $\pm$ 1.93	61.65 $\pm$ 0.61	61.44 $\pm$ 0.87	62.83 $\pm$ 1.01	62.47 $\pm$ 1.14	94.49 $\pm$ 1.56	93.06 $\pm$ 2.20	93.50 $\pm$ 2.01	50.00 $\pm$ 0.00	50.00 $\pm$ 0.00	76.21 $\pm$ 1.08
fraud	93.24 $\pm$ 1.45	93.21 $\pm$ 1.25	94.97 $\pm$ 0.71	95.84 $\pm$ 0.87	93.90 $\pm$ 1.05	94.23 $\pm$ 0.87	94.46 $\pm$ 0.91	93.96 $\pm$ 1.01	94.89 $\pm$ 0.89	66.42 $\pm$ 8.60	50.00 $\pm$ 0.00	76.95 $\pm$ 2.58
glass	84.02 $\pm$ 8.38	94.06 $\pm$ 3.83	90.25 $\pm$ 0.37	99.30 $\pm$ 0.35	99.12 $\pm$ 0.55	95.27 $\pm$ 1.45	93.46 $\pm$ 0.91	93.96 $\pm$ 1.01	94.89 $\pm$ 0.89	66.42 $\pm$ 8.60	50.00 $\pm$ 0.00	76.95 $\pm$ 2.58
hepatitis	99.95 $\pm$ 0.11	98.76 $\pm$ 1.84	99.98 $\pm$ 0.14	98.86 $\pm$ 0.29	99.97 $\pm$ 0.06	99.09 $\pm$ 0.48	99.54 $\pm$ 0.58	98.90 $\pm$ 1.01	98.90 $\pm$ 1.01	99.39 $\pm$ 0.07	50.00 $\pm$ 0.00	89.59 $\pm$ 0.09
laptop	99.96 $\pm$ 0.07	99.97 $\pm$ 0.02	99.98 $\pm$ 0.01	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	99.98 $\pm$ 0.02	99.54 $\pm$ 0.08	99.64 $\pm$ 0.28	99.64 $\pm$ 0.28	50.00 $\pm$ 0.00	50.00 $\pm$ 0.00	49.27 $\pm$ 1.12
inurl	52.16 $\pm$ 0.31	51.88 $\pm$ 0.40	52.28 $\pm$ 0.15	52.35 $\pm$ 0.12	53.06 $\pm$ 0.15	52.53 $\pm$ 0.12	79.02 $\pm$ 1.49	77.71 $\pm$ 0.82	77.23 $\pm$ 0.28	50.00 $\pm$ 0.00	50.00 $\pm$ 0.00	66.79 $\pm$ 0.73
internetads	72.61 $\pm$ 1.19	73.97 $\pm$ 0.44	74.09 $\pm$ 1.75	73.47 $\pm$ 0.45	68.51 $\pm$ 0.74	72.53 $\pm$ 0.15	94.67 $\pm$ 1.52	94.09 $\pm$ 2.42	95.23 $\pm$ 1.47	85.77 $\pm$ 2.62	44.90 $\pm$ 23.57	83.01 $\pm$ 5.81
ionosphere	96.81 $\pm$ 2.22	96.13 $\pm$ 3.45	98.90 $\pm$ 0.41	98.91 $\pm$ 0.31	98.14 $\pm$ 0.79	97.78 $\pm$ 1.23	94.67 $\pm$ 1.52	94.09 $\pm$ 2.42	95.23 $\pm$ 1.47	85.77 $\pm$ 2.62	44.90 $\pm$ 23.57	83.01 $\pm$ 5.81
landsat	65.71 $\pm$ 2.05	60.63 $\pm$ 1.79	65.96 $\pm$ 0.76	67.85 $\pm$ 0.68	61.82 $\pm$ 1.06	64.39 $\pm$ 0.59	51.80 $\pm$ 0.94	50.31 $\pm$ 2.61	48.67 $\pm$ 2.62	50.22 $\pm$ 0.45	50.00 $\pm$ 0.00	50.20 $\pm$ 0.98
letter	48.14 $\pm$ 3.53	42.74 $\pm$ 3.41	40.70 $\pm$ 2.13	40.32 $\pm$ 1.11	47.20 $\pm$ 2.50	43.82 $\pm$ 1.97	38.72 $\pm$ 1.19	37.99 $\pm$ 1.31	37.40 $\pm$ 0.89	50.00 $\pm$ 0.00	50.00 $\pm$ 0.00	42.38 $\pm$ 0.25
lymphography	69.72 $\pm$ 3.02	76.94 $\pm$ 3.25	77.61 $\pm$ 4.70	77.81 $\pm$ 0.88	78.36 $\pm$ 1.34	76.09 $\pm$ 1.17	86.19 $\pm$ 0.59	87.12 $\pm$ 0.35	87.21 $\pm$ 0.73	64.88 $\pm$ 18.23	50.00 $\pm$ 0.00	88.19 $\pm$ 1.78
magic_gamma	79.09 $\pm$ 5.19	85.08 $\pm$ 3.05	79.36 $\pm$ 2.14	80.75 $\pm$ 1.40	77.35 $\pm$ 0.47	80.17 $\pm$ 0.63	82.00 $\pm$ 0.62	86.02 $\pm$ 1.23	87.45 $\pm$ 2.49	85.04 $\pm$ 1.03	50.00 $\pm$ 0.00	75.19 $\pm$ 3.50
maninography	73.53 $\pm$ 1.43	70.05 $\pm$ 3.81	79.13 $\pm$ 1.20	87.05 $\pm$ 1.12	80.10 $\pm$ 0.77	81.37 $\pm$ 0.53	90.21 $\pm$ 0.70	86.02 $\pm$ 1.23	87.45 $\pm$ 2.49	85.04 $\pm$ 1.03	50.00 $\pm$ 0.00	77.55 $\pm$ 0.64
mnist	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	50.00 $\pm$ 0.00	50.00 $\pm$ 0.00	80.00 $\pm$ 0.00
mnist	91.22 $\pm$ 1.88	94.15 $\pm$ 0.97	95.17 $\pm$ 0.85	96.47 $\pm$ 0.63	95.08 $\pm$ 1.06	94.82 $\pm$ 0.49	96.29 $\pm$ 1.29	97.42 $\pm$ 0.83	97.37 $\pm$ 1.38	50.00 $\pm$ 0.00	50.00 $\pm$ 0.00	67.73 $\pm$ 1.71
optdigits	89.21 $\pm$ 1.10	90.77 $\pm$ 1.73	96.45 $\pm$ 2.40	96.17 $\pm$ 2.08	96.28 $\pm$ 0.77	94.00 $\pm$ 2.36	99.96 $\pm$ 0.39	90.29 $\pm$ 0.81	89.23 $\pm$ 0.26	88.80 $\pm$ 0.13	50.00 $\pm$ 0.00	81.66 $\pm$ 0.20
pagoblocks	57.01 $\pm$ 9.78	74.62 $\pm$ 3.27	73.01 $\pm$ 1.77	79.77 $\pm$ 1.92	75.09 $\pm$ 2.65	76.12 $\pm$ 1.04	70.12 $\pm$ 1.89	67.22 $\pm$ 3.15	66.28 $\pm$ 2.78	71.26 $\pm$ 2.38	69.39 $\pm$ 5.15	68.95 $\pm$ 0.96
plma	74.02 $\pm$ 9.67	83.23 $\pm$ 2.89	85.48 $\pm$ 0.32	89.24 $\pm$ 0.37	85.35 $\pm$ 0.48	83.57 $\pm$ 2.48	80.21 $\pm$ 1.18	79.01 $\pm$ 0.76	78.54 $\pm$ 0.62	56.11 $\pm$ 12.22	50.00 $\pm$ 0.00	68.77 $\pm$ 2.26
satellite	94.41 $\pm$ 2.46	93.01 $\pm$ 12.69	99.80 $\pm$ 0.06	99.55 $\pm$ 0.09	99.97 $\pm$ 0.01	96.99 $\pm$ 2.56	99.61 $\pm$ 0.13	99.17 $\pm$ 0.15	98.38 $\pm$ 0.56	69.34 $\pm$ 23.69	50.00 $\pm$ 0.00	83.30 $\pm$ 4.71
satimage-2	99.43 $\pm$ 0.50	99.89 $\pm$ 0.05	99.99 $\pm$ 0.00	99.99 $\pm$ 0.00	99.97 $\pm$ 0.01	99.85 $\pm$ 0.11	99.76 $\pm$ 0.01	99.74 $\pm$ 0.01	99.70 $\pm$ 0.01	99.63 $\pm$ 0.28	50.00 $\pm$ 0.00	89.70 $\pm$ 0.05
shuttle	53.71 $\pm$ 32.40	86.42 $\pm$ 5.89	86.94 $\pm$ 5.22	71.25 $\pm$ 17.43	70.72 $\pm$ 13.69	73.81 $\pm$ 10.45	91.69 $\pm$ 0.28	92.12 $\pm$ 0.27	91.95 $\pm$ 0.22	91.63 $\pm$ 0.28	50.00 $\pm$ 0.00	86.47 $\pm$ 0.96
skin	81.23 $\pm$ 10.57	87.97 $\pm$ 4.57	90.53 $\pm$ 5.25	89.48 $\pm$ 4.52	92.63 $\pm$ 4.06	88.37 $\pm$ 5.41	95.06 $\pm$ 1.20	95.06 $\pm$ 1.20	95.84 $\pm$ 1.23	95.84 $\pm$ 1.22	50.00 $\pm$ 0.00	86.47 $\pm$ 0.96
smtp	79.20 $\pm$ 2.89	79.36 $\pm$ 5.07	83.16 $\pm$ 0.34	83.39 $\pm$ 0.40	78.42 $\pm$ 0.40	80.77 $\pm$ 0.86	82.98 $\pm$ 0.37	83.29 $\pm$ 0.49	83.74 $\pm$ 0.38	50.00 $\pm$ 0.00	50.00 $\pm$ 0.00	43.06 $\pm$ 0.55
spanbase	50.06 $\pm$ 3.03	50.96 $\pm$ 2.48	53.72 $\pm$ 1.46	51.29 $\pm$ 2.33	46.64 $\pm$ 1.97	50.53 $\pm$ 1.10	58.02 $\pm$ 1.49	58.55 $\pm$ 1.07	58.48 $\pm$ 2.57	90.17 $\pm$ 4.76	50.10 $\pm$ 26.50	82.21 $\pm$ 5.22
speech	77.62 $\pm$ 9.15	84.33 $\pm$ 6.20	96.70 $\pm$ 0.74	96.78 $\pm$ 1.19	93.73 $\pm$ 0.80	95.35 $\pm$ 0.07	98.75 $\pm$ 0.07	98.92 $\pm$ 0.02	98.92 $\pm$ 0.04	98.94 $\pm$ 0.05	50.00 $\pm$ 0.00	89.11 $\pm$ 0.02
stamps	94.79 $\pm$ 1.60	95.24 $\pm$ 1.04	96.20 $\pm$ 0.74	96.78 $\pm$ 1.19	93.73 $\pm$ 0.80	95.35 $\pm$ 0.07	98.75 $\pm$ 0.07	98.92 $\pm$ 0.02	98.92 $\pm$ 0.04	56.35 $\pm$ 4.78	48.40 $\pm$ 5.99	59.97 $\pm$ 2.61
thyroid	53.80 $\pm$ 4.97	58.76 $\pm$ 7.63	75.60 $\pm$ 6.18	82.96 $\pm$ 2.86	81.92 $\pm$ 1.93	70.61 $\pm$ 1.11	67.07 $\pm$ 2.73	65.69 $\pm$ 3.50	62.93 $\pm$ 4.76	50.00 $\pm$ 0.00	50.00 $\pm$ 0.00	72.28 $\pm$ 0.42
vertebral	73.59 $\pm$ 6.94	79.11 $\pm$ 2.66	84.59 $\pm$ 3.49	84.81 $\pm$ 3.61	83.99 $\pm$ 1.93	81.22 $\pm$ 1.57	87.25 $\pm$ 1.51	86.66 $\pm$ 0.72	87.40 $\pm$ 0.93	50.00 $\pm$ 0.00	50.00 $\pm$ 0.00	59.07 $\pm$ 0.68
vowels	73.85 $\pm$ 6.07	70.94 $\pm$ 1.78	70.28 $\pm$ 2.82	61.95 $\pm$ 1.35	64.69 $\pm$ 1.66	68.33 $\pm$ 2.15	65.16 $\pm$ 1.34	63.97 $\pm$ 2.75	66.24 $\pm$ 1.45	50.00 $\pm$ 0.00	50.00 $\pm$ 0.00	81.10 $\pm$ 1.77
waveform	98.13 $\pm$ 1.33	99.08 $\pm$ 0.62	99.89 $\pm$ 0.19	99.67 $\pm$ 0.15	99.29 $\pm$ 0.25	98.44 $\pm$ 0.42	98.96 $\pm$ 0.46	98.70 $\pm$ 0.41	97.06 $\pm$ 0.69	50.22 $\pm$ 16.16	45.17 $\pm$ 9.66	78.42 $\pm$ 3.73
wbc	98.06 $\pm$ 2.87	99.05 $\pm$ 0.72	99.51 $\pm$ 0.25	99.07 $\pm$ 0.15	99.81 $\pm$ 0.39	99.79 $\pm$ 0.33	99.91 $\pm$ 0.11	99.74 $\pm$ 0.27	97.06 $\pm$ 0.69	66.04 $\pm$ 26.91	49.98 $\pm$ 0.04	83.09 $\pm$ 5.95
wdbc	52.75 $\pm$ 14.46	54.12 $\pm$ 3.69	47.02 $\pm$ 0.62	47.38 $\pm$ 1.07	46.32 $\pm$ 0.91	72.42 $\pm$ 2.57	84.17 $\pm$ 0.27					

Table 15.1: Average AUPR $\pm$ standard dev. over five seeds for the semi-supervised setting of DTE-NP, k NN baselines with varying hyperparameter (HP) values; $k \in \{5, 10, 20, 40, 50\}$. Also reported is the avg model. We use **bold** and underline respectively to mark the **best** and the worst performance of each model to showcase the variability of performance across different HP settings.

dataset	DTE-NP-5	DTE-NP-10	DTE-NP-20	DTE-NP-40	DTE-NP-50	DTE-NP-avg	KNN-5	KNN-10	KNN-20	KNN-40	KNN-50	KNN-avg
amazon	5.95 $\pm$ 0.00	5.99 $\pm$ 0.00	6.02 $\pm$ 0.00	6.06 $\pm$ 0.00	6.07 $\pm$ 0.00	6.02 $\pm$ 0.00	6.02 $\pm$ 0.00	6.07 $\pm$ 0.00	6.09 $\pm$ 0.00	6.13 $\pm$ 0.00	6.15 $\pm$ 0.00	6.09 $\pm$ 0.00
ala1	11.68 $\pm$ 0.00	11.68 $\pm$ 0.00	11.68 $\pm$ 0.00	11.61 $\pm$ 0.00	11.65 $\pm$ 0.00	11.65 $\pm$ 0.00	11.62 $\pm$ 0.00	11.70 $\pm$ 0.00	11.69 $\pm$ 0.00	11.60 $\pm$ 0.00	11.59 $\pm$ 0.00	11.65 $\pm$ 0.00
amazon	67.49 $\pm$ 0.00	66.73 $\pm$ 0.00	66.04 $\pm$ 0.00	65.39 $\pm$ 0.00	64.87 $\pm$ 0.00	64.87 $\pm$ 0.00	68.07 $\pm$ 0.00	67.90 $\pm$ 0.00	67.26 $\pm$ 0.00	66.27 $\pm$ 0.00	65.79 $\pm$ 0.00	67.06 $\pm$ 0.00
backdoor	55.90 $\pm$ 0.99	47.16 $\pm$ 1.45	38.31 $\pm$ 1.02	31.44 $\pm$ 0.47	20.58 $\pm$ 0.37	40.81 $\pm$ 0.81	46.70 $\pm$ 1.22	37.30 $\pm$ 1.35	29.58 $\pm$ 0.58	24.34 $\pm$ 0.41	22.34 $\pm$ 0.53	32.06 $\pm$ 0.76
breastw	98.51 $\pm$ 0.56	98.19 $\pm$ 0.58	97.56 $\pm$ 0.51	97.13 $\pm$ 0.62	97.05 $\pm$ 0.45	97.69 $\pm$ 0.40	98.97 $\pm$ 0.28	99.01 $\pm$ 0.31	99.08 $\pm$ 0.23	99.15 $\pm$ 0.17	99.16 $\pm$ 0.16	99.08 $\pm$ 0.22
campaign	48.48 $\pm$ 0.00	49.05 $\pm$ 0.00	49.77 $\pm$ 0.00	49.77 $\pm$ 0.00	49.51 $\pm$ 0.00	49.31 $\pm$ 0.00	49.04 $\pm$ 0.00	49.89 $\pm$ 0.00	49.47 $\pm$ 0.00	49.33 $\pm$ 0.00	49.64 $\pm$ 0.00	49.64 $\pm$ 0.00
cardio	76.90 $\pm$ 0.00	77.73 $\pm$ 0.00	78.30 $\pm$ 0.00	79.19 $\pm$ 0.00	79.53 $\pm$ 0.00	78.33 $\pm$ 0.00	77.22 $\pm$ 0.00	78.37 $\pm$ 0.00	79.14 $\pm$ 0.00	80.67 $\pm$ 0.00	79.30 $\pm$ 0.00	79.30 $\pm$ 0.00
cardiologocography	56.55 $\pm$ 0.00	57.18 $\pm$ 0.00	58.19 $\pm$ 0.00	59.42 $\pm$ 0.00	59.95 $\pm$ 0.00	58.26 $\pm$ 0.00	57.43 $\pm$ 0.00	58.37 $\pm$ 0.00	59.44 $\pm$ 0.00	61.41 $\pm$ 0.00	62.19 $\pm$ 0.00	59.77 $\pm$ 0.00
celeba	10.56 $\pm$ 0.44	11.63 $\pm$ 0.49	12.74 $\pm$ 0.52	13.92 $\pm$ 0.58	14.30 $\pm$ 0.59	12.63 $\pm$ 0.51	11.99 $\pm$ 0.57	13.22 $\pm$ 0.61	14.50 $\pm$ 0.58	15.70 $\pm$ 0.65	16.10 $\pm$ 0.68	14.31 $\pm$ 0.60
cover	21.14 $\pm$ 0.39	21.38 $\pm$ 0.54	21.16 $\pm$ 0.43	20.67 $\pm$ 0.41	20.52 $\pm$ 0.42	20.97 $\pm$ 0.43	21.36 $\pm$ 0.76	21.22 $\pm$ 0.69	20.59 $\pm$ 0.33	20.00 $\pm$ 0.42	19.94 $\pm$ 0.44	20.62 $\pm$ 0.44
cover	63.67 $\pm$ 5.21	57.85 $\pm$ 3.52	51.55 $\pm$ 3.10	44.58 $\pm$ 2.49	42.11 $\pm$ 2.29	51.95 $\pm$ 2.90	55.15 $\pm$ 3.45	48.67 $\pm$ 2.84	41.44 $\pm$ 2.04	33.72 $\pm$ 1.51	31.69 $\pm$ 1.35	42.14 $\pm$ 2.20
donors	93.23 $\pm$ 0.80	91.25 $\pm$ 0.72	88.17 $\pm$ 0.95	83.92 $\pm$ 1.29	82.31 $\pm$ 1.32	87.78 $\pm$ 0.99	89.44 $\pm$ 0.96	85.33 $\pm$ 1.15	80.15 $\pm$ 1.33	73.68 $\pm$ 1.32	71.00 $\pm$ 1.30	79.92 $\pm$ 1.13
fault	62.03 $\pm$ 0.00	61.58 $\pm$ 0.00	61.29 $\pm$ 0.00	61.98 $\pm$ 0.00	62.31 $\pm$ 0.00	61.84 $\pm$ 0.00	61.98 $\pm$ 0.00	61.16 $\pm$ 0.00	61.92 $\pm$ 0.00	63.67 $\pm$ 0.00	64.06 $\pm$ 0.00	62.56 $\pm$ 0.00
fraud	40.60 $\pm$ 6.67	43.77 $\pm$ 5.38	43.03 $\pm$ 4.92	39.91 $\pm$ 4.75	38.80 $\pm$ 4.94	41.22 $\pm$ 5.02	42.35 $\pm$ 5.61	44.96 $\pm$ 3.90	41.19 $\pm$ 3.73	37.33 $\pm$ 4.15	36.42 $\pm$ 4.12	40.45 $\pm$ 4.07
glass	60.15 $\pm$ 6.89	47.75 $\pm$ 5.62	37.27 $\pm$ 4.92	31.23 $\pm$ 3.12	30.48 $\pm$ 3.45	41.38 $\pm$ 4.46	44.05 $\pm$ 6.38	32.96 $\pm$ 3.74	29.87 $\pm$ 3.75	26.62 $\pm$ 2.65	26.04 $\pm$ 3.61	31.91 $\pm$ 3.73
hepatitis	99.47 $\pm$ 0.73	97.98 $\pm$ 1.43	91.71 $\pm$ 2.26	81.65 $\pm$ 3.68	78.38 $\pm$ 3.85	89.95 $\pm$ 1.80	91.10 $\pm$ 4.41	82.95 $\pm$ 3.16	64.12 $\pm$ 5.59	64.29 $\pm$ 5.08	64.33 $\pm$ 6.16	70.62 $\pm$ 4.48
lftp	98.52 $\pm$ 0.37	95.38 $\pm$ 2.26	88.66 $\pm$ 1.01	84.43 $\pm$ 2.65	80.40 $\pm$ 4.55	89.48 $\pm$ 1.90	100.00 $\pm$ 0.00	98.01 $\pm$ 3.98	91.24 $\pm$ 1.42	91.44 $\pm$ 1.25	91.28 $\pm$ 1.36	94.39 $\pm$ 1.31
lftp	91.11 $\pm$ 0.00	9.09 $\pm$ 0.00	9.07 $\pm$ 0.00	9.06 $\pm$ 0.00	9.06 $\pm$ 0.00	9.08 $\pm$ 0.00	8.92 $\pm$ 0.00	8.94 $\pm$ 0.00	8.99 $\pm$ 0.00	8.98 $\pm$ 0.00	8.99 $\pm$ 0.00	8.96 $\pm$ 0.00
interpreds	52.20 $\pm$ 0.48	49.76 $\pm$ 0.00	48.19 $\pm$ 0.00	47.56 $\pm$ 0.00	47.45 $\pm$ 0.00	49.03 $\pm$ 0.00	49.22 $\pm$ 0.40	47.20 $\pm$ 0.00	46.93 $\pm$ 0.00	46.95 $\pm$ 0.00	46.94 $\pm$ 0.00	47.47 $\pm$ 0.00
ionosphere	98.72 $\pm$ 0.48	98.46 $\pm$ 0.54	98.27 $\pm$ 0.42	97.44 $\pm$ 0.50	96.93 $\pm$ 0.61	97.96 $\pm$ 0.46	97.86 $\pm$ 0.60	98.11 $\pm$ 0.52	97.04 $\pm$ 0.60	94.12 $\pm$ 1.45	92.90 $\pm$ 1.65	96.01 $\pm$ 0.78
ionosphere	56.14 $\pm$ 0.00	54.25 $\pm$ 0.00	50.75 $\pm$ 0.00	46.43 $\pm$ 0.00	45.17 $\pm$ 0.00	50.55 $\pm$ 0.00	54.85 $\pm$ 0.00	50.62 $\pm$ 0.00	45.18 $\pm$ 0.00	41.32 $\pm$ 0.00	40.50 $\pm$ 0.00	46.19 $\pm$ 0.00
letter	8.86 $\pm$ 0.00	8.78 $\pm$ 0.00	8.67 $\pm$ 0.00	8.54 $\pm$ 0.00	8.50 $\pm$ 0.00	8.67 $\pm$ 0.00	8.70 $\pm$ 0.00	8.58 $\pm$ 0.00	8.41 $\pm$ 0.00	8.27 $\pm$ 0.00	8.22 $\pm$ 0.00	8.44 $\pm$ 0.00
lymphography	97.27 $\pm$ 5.45	96.07 $\pm$ 6.70	96.07 $\pm$ 6.70	95.68 $\pm$ 6.60	96.16 $\pm$ 6.43	95.68 $\pm$ 6.60	98.61 $\pm$ 1.02	98.43 $\pm$ 0.52	98.13 $\pm$ 0.52	98.70 $\pm$ 0.53	98.70 $\pm$ 0.53	98.57 $\pm$ 0.65
magic gamma	86.20 $\pm$ 0.00	85.86 $\pm$ 0.00	85.28 $\pm$ 0.00	84.29 $\pm$ 0.00	84.29 $\pm$ 0.00	85.28 $\pm$ 0.00	85.86 $\pm$ 0.00	85.45 $\pm$ 0.00	84.51 $\pm$ 0.00	83.61 $\pm$ 0.00	83.25 $\pm$ 0.00	84.50 $\pm$ 0.00
maninography	42.14 $\pm$ 0.00	41.51 $\pm$ 0.00	40.67 $\pm$ 0.00	40.27 $\pm$ 0.00	40.27 $\pm$ 0.00	41.04 $\pm$ 0.00	41.27 $\pm$ 0.00	40.55 $\pm$ 0.00	40.24 $\pm$ 0.00	38.97 $\pm$ 0.00	38.10 $\pm$ 0.00	39.89 $\pm$ 0.00
mnist	74.43 $\pm$ 0.00	73.09 $\pm$ 0.00	71.84 $\pm$ 0.00	70.69 $\pm$ 0.00	70.36 $\pm$ 0.00	72.08 $\pm$ 0.00	72.72 $\pm$ 0.00	71.40 $\pm$ 0.00	70.09 $\pm$ 0.00	69.02 $\pm$ 0.00	68.00 $\pm$ 0.00	70.36 $\pm$ 0.00
mnist	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00
mnist	34.44 $\pm$ 0.00	30.53 $\pm$ 0.00	26.28 $\pm$ 0.00	22.67 $\pm$ 0.00	21.01 $\pm$ 0.00	27.11 $\pm$ 0.00	29.11 $\pm$ 0.00	27.77 $\pm$ 0.00	27.10 $\pm$ 0.00	26.77 $\pm$ 0.00	26.13 $\pm$ 0.00	27.10 $\pm$ 0.00
mnist	62.78 $\pm$ 0.00	62.52 $\pm$ 0.00	62.20 $\pm$ 0.00	61.02 $\pm$ 0.00	60.30 $\pm$ 0.00	61.76 $\pm$ 0.00	67.60 $\pm$ 0.00	67.77 $\pm$ 0.00	67.60 $\pm$ 0.00	67.87 $\pm$ 0.00	67.87 $\pm$ 0.00	67.87 $\pm$ 0.00
mnist	97.68 $\pm$ 0.00	97.31 $\pm$ 0.00	96.28 $\pm$ 0.00	90.01 $\pm$ 0.00	86.06 $\pm$ 0.00	93.39 $\pm$ 0.00	96.99 $\pm$ 0.00	95.65 $\pm$ 0.00	91.40 $\pm$ 0.00	70.28 $\pm$ 0.00	67.39 $\pm$ 0.00	82.34 $\pm$ 0.00
mnist	80.27 $\pm$ 1.65	78.05 $\pm$ 2.13	75.87 $\pm$ 2.40	74.73 $\pm$ 2.71	74.49 $\pm$ 2.71	76.08 $\pm$ 2.22	75.66 $\pm$ 2.91	73.62 $\pm$ 2.39	73.42 $\pm$ 2.98	73.71 $\pm$ 2.79	73.63 $\pm$ 2.86	74.01 $\pm$ 2.75
mnist	85.98 $\pm$ 0.00	85.74 $\pm$ 0.00	85.17 $\pm$ 0.00	84.15 $\pm$ 0.00	83.72 $\pm$ 0.00	84.35 $\pm$ 0.00	86.01 $\pm$ 0.00	85.31 $\pm$ 0.00	84.02 $\pm$ 0.00	82.19 $\pm$ 0.00	81.56 $\pm$ 0.00	83.82 $\pm$ 0.00
mnist	99.16 $\pm$ 0.00	96.64 $\pm$ 0.00	97.02 $\pm$ 0.00	97.39 $\pm$ 0.00	97.42 $\pm$ 0.00	96.92 $\pm$ 0.00	97.86 $\pm$ 0.00	97.21 $\pm$ 0.00	97.28 $\pm$ 0.00	97.22 $\pm$ 0.00	97.22 $\pm$ 0.00	97.22 $\pm$ 0.00
mnist	98.92 $\pm$ 0.23	98.31 $\pm$ 0.20	96.81 $\pm$ 0.31	94.52 $\pm$ 0.43	93.30 $\pm$ 0.50	96.37 $\pm$ 0.25	98.31 $\pm$ 0.34	96.30 $\pm$ 0.30	94.09 $\pm$ 0.51	88.39 $\pm$ 0.08	86.43 $\pm$ 0.46	97.38 $\pm$ 0.16
mnist	56.70 $\pm$ 7.16	54.77 $\pm$ 7.80	54.75 $\pm$ 7.81	54.76 $\pm$ 7.81	48.71 $\pm$ 10.23	53.94 $\pm$ 5.73	50.26 $\pm$ 5.73	50.20 $\pm$ 5.74	50.18 $\pm$ 5.75	50.33 $\pm$ 5.85	50.41 $\pm$ 5.72	50.27 $\pm$ 5.76
mnist	83.93 $\pm$ 0.00	83.42 $\pm$ 0.00	83.03 $\pm$ 0.00	82.73 $\pm$ 0.00	82.63 $\pm$ 0.00	83.15 $\pm$ 0.00	83.32 $\pm$ 0.00	82.70 $\pm$ 0.00	82.41 $\pm$ 0.00	82.17 $\pm$ 0.00	82.11 $\pm$ 0.00	82.54 $\pm$ 0.00
mnist	3.02 $\pm$ 0.00	2.89 $\pm$ 0.00	2.70 $\pm$ 0.00	2.76 $\pm$ 0.00	2.82 $\pm$ 0.00	2.82 $\pm$ 0.00	2.80 $\pm$ 0.00	2.73 $\pm$ 0.00	2.74 $\pm$ 0.00	2.76 $\pm$ 0.00	2.74 $\pm$ 0.00	2.75 $\pm$ 0.00
mnist	82.50 $\pm$ 3.71	77.11 $\pm$ 4.30	69.99 $\pm$ 5.29	65.85 $\pm$ 6.16	64.61 $\pm$ 6.16	72.02 $\pm$ 4.82	73.26 $\pm$ 7.70	65.58 $\pm$ 7.71	63.12 $\pm$ 6.69	62.09 $\pm$ 7.20	61.57 $\pm$ 7.18	65.12 $\pm$ 7.10
mnist	77.22 $\pm$ 0.00	77.53 $\pm$ 0.00	77.26 $\pm$ 0.00	76.43 $\pm$ 0.00	74.75 $\pm$ 0.00	76.64 $\pm$ 0.00	80.91 $\pm$ 0.00	81.09 $\pm$ 0.00	81.90 $\pm$ 0.00	81.90 $\pm$ 0.00	81.93 $\pm$ 0.00	81.47 $\pm$ 0.00
mnist	43.77 $\pm$ 5.50	31.72 $\pm$ 3.49	24.99 $\pm$ 2.97	21.10 $\pm$ 2.37	20.29 $\pm$ 2.40	28.38 $\pm$ 3.31	25.07 $\pm$ 3.19	21.55 $\pm$ 2.53	19.65 $\pm$ 2.31	18.09 $\pm$ 1.91	17.76 $\pm$ 2.02	20.42 $\pm$ 2.34
mnist	31.68 $\pm$ 0.00	30.30 $\pm$ 0.00	29.54 $\pm$ 0.00	27.85 $\pm$ 0.00	27.32 $\pm$ 0.00	29.34 $\pm$ 0.00	30.21 $\pm$ 0.00	28.75 $\pm$ 0.00	27.41 $\pm$ 0.00	24.27 $\pm$ 0.00	22.44 $\pm$ 0.00	26.62 $\pm$ 0.00
mnist	26.96 $\pm$ 0.00	26.71 $\pm$ 0.00	25.68 $\pm$ 0.00	24.67 $\pm$ 0.00	24.30 $\pm$ 0.00	25.70 $\pm$ 0.00	27.00 $\pm$ 0.00	25.82 $\pm$ 0.00	23.77 $\pm$ 0.00	21.13 $\pm$ 0.00	20.92 $\pm$ 0.00	24.92 $\pm$ 0.00
mnist	96.59 $\pm$ 2.18	93.20 $\pm$ 4.25	90.32 $\pm$ 6.04	90.14 $\pm$ 5.60	88.96 $\pm$ 6.03	91.84 $\pm$ 5.57	89.48 $\pm$ 5.58	89.07 $\pm$ 5.41	82.05 $\pm$ 5.59	82.37 $\pm$ 5.91	91.82 $\pm$ 3.32	90.15 $\pm$ 3.97
mnist	92.08 $\pm$ 6.52	89.03 $\pm$ 6.77	86.03 $\pm$ 5.81	83.79 $\pm$ 5.59	83.41 $\pm$ 5.19	86.87 $\pm$ 5.84	85.35 $\pm$ 5.46	83.75 $\pm$ 5.56	82.05 $\pm$ 5.59	82.37 $\pm$ 5.91	82.33 $\pm$ 5.93	83.17 $\pm$ 5.03
mnist	13.43 $\pm$											

Table 15.2: Average AUPR $\pm$ standard dev. over five seeds for the semi-supervised setting of ICL and DTE-C baselines with varying hyperparameter (HP) values; For ICL, the learning rate $\in \{0.1, 0.02, 0.001, 0.0001, 1e-05\}$, for DTE-C, $k \in \{5, 10, 20, 40, 50\}$. Also reported is the avg model. We use **bold** and underline respectively to mark the **best** and the **worst** performance of each model to showcase the variability of performance across different HP settings.

dataset	ICL-0.1	ICL-0.01	ICL-0.001	ICL-0.0001	ICL-1e-05	ICL-avg	DTE-C-5	DTE-C-10	DTE-C-20	DTE-C-40	DTE-C-50	DTE-C-avg
abi	5.50 $\pm$ 0.09	5.30 $\pm$ 0.07	5.50 $\pm$ 0.08	5.59 $\pm$ 0.04	5.46 $\pm$ 0.01	5.40 $\pm$ 0.02	5.76 $\pm$ 0.03	5.82 $\pm$ 0.02	5.72 $\pm$ 0.05	5.73 $\pm$ 0.05	5.91 $\pm$ 0.00	5.79 $\pm$ 0.01
amazon	10.06 $\pm$ 0.10	10.08 $\pm$ 0.03	9.91 $\pm$ 0.13	10.01 $\pm$ 0.05	9.99 $\pm$ 0.02	10.01 $\pm$ 0.05	11.01 $\pm$ 0.43	10.99 $\pm$ 1.11	11.05 $\pm$ 0.33	9.52 $\pm$ 0.00	10.42 $\pm$ 0.40	10.42 $\pm$ 0.40
amthyroid	58.53 $\pm$ 13.33	39.69 $\pm$ 5.23	53.66 $\pm$ 3.08	53.85 $\pm$ 5.44	55.94 $\pm$ 2.70	52.34 $\pm$ 3.72	82.46 $\pm$ 0.50	81.52 $\pm$ 1.18	83.27 $\pm$ 0.63	81.52 $\pm$ 1.18	13.81 $\pm$ 0.00	68.86 $\pm$ 0.20
backdoor	85.81 $\pm$ 2.45	88.14 $\pm$ 1.38	88.99 $\pm$ 1.12	88.96 $\pm$ 1.20	88.96 $\pm$ 1.11	87.68 $\pm$ 1.03	43.90 $\pm$ 3.90	61.21 $\pm$ 2.58	63.64 $\pm$ 1.61	48.3 $\pm$ 0.09	35.68 $\pm$ 0.79	35.68 $\pm$ 0.79
breastw	98.65 $\pm$ 0.81	98.66 $\pm$ 0.51	98.98 $\pm$ 0.57	98.50 $\pm$ 0.42	95.32 $\pm$ 1.11	97.98 $\pm$ 0.31	89.69 $\pm$ 2.54	90.04 $\pm$ 1.09	98.56 $\pm$ 0.68	99.22 $\pm$ 0.11	94.00 $\pm$ 0.69	94.00 $\pm$ 0.69
campaign	47.46 $\pm$ 0.41	44.03 $\pm$ 2.05	47.94 $\pm$ 0.25	49.22 $\pm$ 0.91	51.41 $\pm$ 0.51	48.01 $\pm$ 0.47	49.90 $\pm$ 2.01	46.77 $\pm$ 1.41	48.40 $\pm$ 1.60	20.25 $\pm$ 0.00	37.11 $\pm$ 0.64	37.11 $\pm$ 0.64
cardio	48.81 $\pm$ 14.30	65.70 $\pm$ 11.26	63.55 $\pm$ 5.07	61.75 $\pm$ 2.55	29.67 $\pm$ 1.62	53.93 $\pm$ 3.02	69.32 $\pm$ 0.43	70.19 $\pm$ 0.63	70.25 $\pm$ 0.47	17.78 $\pm$ 0.59	17.55 $\pm$ 0.00	49.02 $\pm$ 0.29
cardiolography	43.21 $\pm$ 5.03	52.56 $\pm$ 1.14	50.69 $\pm$ 2.12	49.01 $\pm$ 1.29	39.79 $\pm$ 1.23	47.95 $\pm$ 1.34	53.83 $\pm$ 0.94	53.56 $\pm$ 0.73	49.12 $\pm$ 1.17	36.12 $\pm$ 0.00	45.75 $\pm$ 0.91	45.75 $\pm$ 0.91
celeba	12.89 $\pm$ 1.17	13.90 $\pm$ 1.19	13.09 $\pm$ 1.73	13.50 $\pm$ 1.15	12.97 $\pm$ 0.55	13.27 $\pm$ 0.47	15.16 $\pm$ 1.05	13.87 $\pm$ 0.59	12.60 $\pm$ 1.22	4.31 $\pm$ 0.09	4.31 $\pm$ 0.09	10.05 $\pm$ 0.49
census	20.29 $\pm$ 1.27	20.38 $\pm$ 2.24	23.32 $\pm$ 0.50	23.68 $\pm$ 0.70	22.37 $\pm$ 0.58	22.01 $\pm$ 0.42	18.39 $\pm$ 0.83	17.28 $\pm$ 0.52	17.45 $\pm$ 0.64	1.66 $\pm$ 0.20	11.66 $\pm$ 0.20	15.29 $\pm$ 0.31
cover	22.58 $\pm$ 13.63	10.32 $\pm$ 4.03	35.90 $\pm$ 22.12	47.50 $\pm$ 19.76	40.81 $\pm$ 22.18	31.42 $\pm$ 3.97	71.28 $\pm$ 4.54	61.74 $\pm$ 4.63	38.26 $\pm$ 3.29	1.94 $\pm$ 0.07	1.94 $\pm$ 0.07	35.03 $\pm$ 0.82
donors	39.48 $\pm$ 10.00	77.81 $\pm$ 8.28	82.59 $\pm$ 7.03	91.60 $\pm$ 4.15	92.24 $\pm$ 2.17	76.75 $\pm$ 3.27	76.07 $\pm$ 1.84	62.97 $\pm$ 0.60	53.74 $\pm$ 2.07	18.86 $\pm$ 15.32	11.20 $\pm$ 0.14	45.30 $\pm$ 2.85
fault	65.37 $\pm$ 3.40	64.58 $\pm$ 1.81	63.83 $\pm$ 0.74	63.93 $\pm$ 0.94	64.32 $\pm$ 0.96	64.01 $\pm$ 1.09	63.92 $\pm$ 1.40	66.62 $\pm$ 0.60	63.62 $\pm$ 1.06	51.69 $\pm$ 0.25	51.74 $\pm$ 0.32	58.79 $\pm$ 0.41
fraud	51.45 $\pm$ 12.34	49.97 $\pm$ 9.27	56.24 $\pm$ 9.47	62.41 $\pm$ 10.30	72.24 $\pm$ 5.48	58.46 $\pm$ 7.58	68.85 $\pm$ 10.35	57.17 $\pm$ 4.08	24.97 $\pm$ 21.19	0.34 $\pm$ 0.03	0.34 $\pm$ 0.03	30.33 $\pm$ 4.25
glass	49.20 $\pm$ 15.55	69.98 $\pm$ 13.15	88.19 $\pm$ 6.00	87.25 $\pm$ 7.44	87.79 $\pm$ 8.75	76.86 $\pm$ 1.38	68.85 $\pm$ 10.35	57.17 $\pm$ 4.08	24.97 $\pm$ 21.19	0.34 $\pm$ 0.03	0.34 $\pm$ 0.03	27.78 $\pm$ 2.62
hepatitis	90.85 $\pm$ 0.30	92.89 $\pm$ 3.11	90.79 $\pm$ 4.01	98.94 $\pm$ 2.13	99.93 $\pm$ 0.14	99.88 $\pm$ 1.76	98.64 $\pm$ 1.76	95.49 $\pm$ 1.81	96.17 $\pm$ 4.38	54.16 $\pm$ 11.62	27.45 $\pm$ 1.71	74.38 $\pm$ 1.79
htrp	94.85 $\pm$ 0.93	95.76 $\pm$ 3.25	97.86 $\pm$ 1.72	98.94 $\pm$ 0.25	99.55 $\pm$ 0.61	97.96 $\pm$ 2.20	56.93 $\pm$ 7.39	60.07 $\pm$ 7.08	48.53 $\pm$ 7.22	0.24 $\pm$ 0.00	0.24 $\pm$ 0.00	45.08 $\pm$ 0.07
imdb	10.09 $\pm$ 0.00	10.02 $\pm$ 0.07	10.16 $\pm$ 0.04	10.18 $\pm$ 0.05	10.41 $\pm$ 0.04	10.17 $\pm$ 0.03	8.21 $\pm$ 0.63	9.04 $\pm$ 0.47	8.79 $\pm$ 0.30	9.52 $\pm$ 0.00	9.52 $\pm$ 0.00	9.21 $\pm$ 0.02
internetads	57.32 $\pm$ 2.19	60.94 $\pm$ 1.14	62.86 $\pm$ 1.88	63.83 $\pm$ 0.88	47.64 $\pm$ 2.42	58.92 $\pm$ 1.04	60.45 $\pm$ 4.51	56.24 $\pm$ 7.74	57.78 $\pm$ 0.93	31.53 $\pm$ 0.00	31.53 $\pm$ 0.00	47.51 $\pm$ 1.92
ionosphere	96.94 $\pm$ 2.27	99.00 $\pm$ 0.36	98.91 $\pm$ 0.48	97.54 $\pm$ 1.39	97.54 $\pm$ 1.39	97.92 $\pm$ 1.11	96.52 $\pm$ 0.77	96.49 $\pm$ 1.17	96.68 $\pm$ 0.68	86.70 $\pm$ 4.57	56.47 $\pm$ 15.94	86.53 $\pm$ 4.28
landsat	58.66 $\pm$ 1.95	56.09 $\pm$ 1.07	55.72 $\pm$ 1.20	55.69 $\pm$ 0.77	57.23 $\pm$ 0.40	55.98 $\pm$ 0.60	36.22 $\pm$ 0.76	36.06 $\pm$ 2.33	34.97 $\pm$ 0.09	34.32 $\pm$ 0.00	10.09 $\pm$ 0.04	34.98 $\pm$ 0.81
letter	11.92 $\pm$ 1.58	10.84 $\pm$ 0.96	9.41 $\pm$ 0.34	9.71 $\pm$ 0.55	15.87 $\pm$ 1.71	11.44 $\pm$ 0.43	8.92 $\pm$ 0.10	9.01 $\pm$ 0.23	8.96 $\pm$ 0.09	11.76 $\pm$ 0.00	11.76 $\pm$ 0.00	10.09 $\pm$ 0.04
limbography	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	72.08 $\pm$ 7.20	93.30 $\pm$ 2.48	83.45 $\pm$ 12.01	72.50 $\pm$ 16.37	7.89 $\pm$ 0.50	68.82 $\pm$ 5.96
magic-gamma	73.92 $\pm$ 3.66	81.92 $\pm$ 2.83	82.64 $\pm$ 0.71	82.47 $\pm$ 0.98	83.45 $\pm$ 0.84	80.71 $\pm$ 1.03	80.01 $\pm$ 0.45	80.46 $\pm$ 0.29	80.15 $\pm$ 0.44	66.79 $\pm$ 18.08	52.03 $\pm$ 0.00	77.29 $\pm$ 3.52
manography	33.51 $\pm$ 8.47	37.72 $\pm$ 7.54	27.06 $\pm$ 2.69	26.94 $\pm$ 1.92	12.30 $\pm$ 1.13	28.51 $\pm$ 2.27	35.09 $\pm$ 3.35	37.47 $\pm$ 2.68	32.92 $\pm$ 4.57	35.30 $\pm$ 2.06	4.23 $\pm$ 0.00	28.91 $\pm$ 1.92
mnist	39.27 $\pm$ 1.38	50.40 $\pm$ 3.02	54.46 $\pm$ 1.34	63.49 $\pm$ 1.45	65.20 $\pm$ 1.13	41.08 $\pm$ 1.34	18.29 $\pm$ 1.20	35.17 $\pm$ 5.05	12.82 $\pm$ 3.07	5.11 $\pm$ 0.00	16.80 $\pm$ 0.00	62.46 $\pm$ 0.00
mnist	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	64.41 $\pm$ 1.43	68.08 $\pm$ 1.72	63.78 $\pm$ 1.40	5.11 $\pm$ 0.00	6.14 $\pm$ 0.00	10.77 $\pm$ 1.02
mnist	27.13 $\pm$ 3.47	36.36 $\pm$ 3.03	63.86 $\pm$ 2.44	61.20 $\pm$ 0.68	61.20 $\pm$ 0.68	64.90 $\pm$ 1.13	11.59 $\pm$ 3.65	45.92 $\pm$ 3.65	40.53 $\pm$ 1.47	5.11 $\pm$ 0.00	6.14 $\pm$ 0.00	10.77 $\pm$ 1.02
mnist	64.80 $\pm$ 3.36	66.26 $\pm$ 4.89	63.38 $\pm$ 3.59	73.08 $\pm$ 7.23	70.36 $\pm$ 5.65	58.53 $\pm$ 6.17	53.63 $\pm$ 4.87	45.92 $\pm$ 3.65	40.53 $\pm$ 1.47	5.11 $\pm$ 0.00	6.14 $\pm$ 0.00	10.77 $\pm$ 1.02
mnist	39.01 $\pm$ 20.69	66.20 $\pm$ 12.32	63.38 $\pm$ 3.59	73.08 $\pm$ 7.23	70.36 $\pm$ 5.65	58.53 $\pm$ 6.17	53.63 $\pm$ 4.87	45.92 $\pm$ 3.65	40.53 $\pm$ 1.47	5.11 $\pm$ 0.00	6.14 $\pm$ 0.00	10.77 $\pm$ 1.02
mnist	74.29 $\pm$ 3.85	71.41 $\pm$ 2.85	68.43 $\pm$ 0.25	90.38 $\pm$ 0.14	86.65 $\pm$ 0.69	75.31 $\pm$ 2.15	67.82 $\pm$ 2.03	65.33 $\pm$ 3.71	60.04 $\pm$ 5.31	55.55 $\pm$ 11.96	44.44 $\pm$ 0.00	54.53 $\pm$ 1.01
mnist	76.70 $\pm$ 8.77	86.99 $\pm$ 1.74	88.43 $\pm$ 0.25	90.38 $\pm$ 0.14	86.65 $\pm$ 0.69	75.31 $\pm$ 2.15	67.82 $\pm$ 2.03	65.33 $\pm$ 3.71	60.04 $\pm$ 5.31	55.55 $\pm$ 11.96	44.44 $\pm$ 0.00	54.53 $\pm$ 1.01
mnist	37.08 $\pm$ 8.31	77.33 $\pm$ 3.36	96.77 $\pm$ 0.59	95.65 $\pm$ 0.53	90.77 $\pm$ 13.14	71.96 $\pm$ 7.98	83.45 $\pm$ 5.57	82.80 $\pm$ 5.31	81.25 $\pm$ 5.02	24.96 $\pm$ 27.32	2.22 $\pm$ 0.00	43.90 $\pm$ 0.54
mnist	37.77 $\pm$ 1.43	99.09 $\pm$ 0.23	99.82 $\pm$ 0.12	99.91 $\pm$ 0.05	99.72 $\pm$ 0.12	99.28 $\pm$ 0.28	94.26 $\pm$ 0.06	94.08 $\pm$ 0.16	93.19 $\pm$ 0.08	88.77 $\pm$ 1.78	13.33 $\pm$ 0.00	76.73 $\pm$ 0.36
mnist	44.77 $\pm$ 20.35	73.00 $\pm$ 9.87	68.01 $\pm$ 9.89	49.36 $\pm$ 11.31	48.40 $\pm$ 11.06	56.71 $\pm$ 7.09	68.88 $\pm$ 0.57	70.08 $\pm$ 0.63	69.80 $\pm$ 0.53	69.53 $\pm$ 0.66	31.42 $\pm$ 0.19	62.54 $\pm$ 0.43
mnist	38.06 $\pm$ 20.67	42.88 $\pm$ 7.84	38.77 $\pm$ 14.99	37.40 $\pm$ 0.86	36.22 $\pm$ 4.31	38.07 $\pm$ 3.38	50.08 $\pm$ 5.85	52.14 $\pm$ 8.82	41.15 $\pm$ 7.61	15.36 $\pm$ 9.94	0.07 $\pm$ 0.01	31.76 $\pm$ 4.31
mnist	80.80 $\pm$ 1.90	81.18 $\pm$ 2.88	85.16 $\pm$ 0.70	85.42 $\pm$ 0.58	82.51 $\pm$ 0.72	83.01 $\pm$ 0.50	89.44 $\pm$ 0.47	83.59 $\pm$ 0.51	84.12 $\pm$ 0.39	57.05 $\pm$ 0.00	57.05 $\pm$ 0.00	73.05 $\pm$ 0.13
mnist	3.74 $\pm$ 1.02	3.79 $\pm$ 0.38	3.92 $\pm$ 0.39	3.46 $\pm$ 0.20	3.41 $\pm$ 0.23	3.66 $\pm$ 0.23	2.78 $\pm$ 0.12	3.32 $\pm$ 0.58	2.85 $\pm$ 0.16	3.26 $\pm$ 0.00	3.26 $\pm$ 0.00	3.09 $\pm$ 0.15
mnist	45.19 $\pm$ 8.30	52.05 $\pm$ 12.46	80.41 $\pm$ 3.57	80.49 $\pm$ 4.25	74.72 $\pm$ 8.62	66.57 $\pm$ 4.58	61.33 $\pm$ 7.30	57.64 $\pm$ 7.65	53.34 $\pm$ 6.09	56.82 $\pm$ 10.61	24.85 $\pm$ 19.23	50.80 $\pm$ 7.62
mnist	72.79 $\pm$ 3.97	56.66 $\pm$ 5.89	56.38 $\pm$ 5.72	60.13 $\pm$ 6.20	30.60 $\pm$ 3.56	55.29 $\pm$ 1.31	80.34 $\pm$ 1.28	81.93 $\pm$ 2.54	82.86 $\pm$ 1.27	83.14 $\pm$ 1.08	4.81 $\pm$ 0.00	66.61 $\pm$ 0.60
mnist	25.70 $\pm$ 4.65	31.37 $\pm$ 8.07	50.83 $\pm$ 13.40	59.35 $\pm$ 3.61	54.90 $\pm$ 6.08	44.44 $\pm$ 4.60	34.31 $\pm$ 5.07	31.45 $\pm$ 6.44	28.95 $\pm$ 5.65	25.00 $\pm$ 4.75	21.17 $\pm$ 4.01	28.18 $\pm$ 4.88
mnist	21.00 $\pm$ 15.67	21.67 $\pm$ 7.57	28.45 $\pm$ 5.52	28.35 $\pm$ 6.64	22.16 $\pm$ 3.78	28.91 $\pm$ 3.68	39.58 $\pm$ 4.05	41.45 $\pm$ 3.27	44.91 $\pm$ 4.07	5.65 $\pm$ 0.00	5.65 $\pm$ 0.00	27.84 $\pm$ 1.66
mnist	51.44 $\pm$ 2.31	34.05 $\pm$ 18.32	28.41 $\pm$ 9.48	9.51 $\pm$ 0.97	20.22 $\pm$ 4.30	28.91 $\pm$ 3.68	39.58 $\pm$ 4.05	41.45 $\pm$ 3.27	44.91 $\pm$ 4.07	5.65 $\pm$ 0.00	5.65 $\pm$ 0.00	8.24 $\pm$ 0.18
mnist	82.08 $\pm$ 11.29	91.06 $\pm$ 5.01	99.04 $\pm$ 1.73	99.16 $\pm$ 1.35	96.58 $\pm$ 2.18	93.88 $\pm$ 3.40	36.71 $\pm$ 3.31	33.53 $\pm$ 4.49	40.09 $\pm$ 7.65	67.56 $\pm$ 7.06	8.72 $\pm$ 1.13	37.32 $\pm$ 1.43
mnist	66.53 $\pm$ 18.38	85.94 $\$										

Table 16.1: Average F1 score $\pm$ standard dev. over five seeds for the semi-supervised setting of DTE-NP, k NN baselines with varying hyperparameter (HP) values; $k \in \{5, 10, 20, 40, 50\}$. Also reported is the avg model. We use **bold** and underline respectively to mark the **best** and the worst performance of each model to showcase the variability of performance across different HP settings.

dataset	DTE-NP-5	DTE-NP-10	DTE-NP-20	DTE-NP-40	DTE-NP-50	DTE-NP-avg	KNN-5	KNN-10	KNN-20	KNN-40	KNN-50	KNN-avg
aloi	5.90 $\pm$ 0.00	5.84 $\pm$ 0.00	<u>5.70</u> $\pm$ 0.00	5.90 $\pm$ 0.00	5.97 $\pm$ 0.00	5.86 $\pm$ 0.00	5.90 $\pm$ 0.00	5.64 $\pm$ 0.00	5.97 $\pm$ 0.00	6.17 $\pm$ 0.00	6.37 $\pm$ 0.00	6.01 $\pm$ 0.00
amazon	10.80 $\pm$ 0.00	10.80 $\pm$ 0.00	10.20 $\pm$ 0.00	10.20 $\pm$ 0.00	10.60 $\pm$ 0.00	10.20 $\pm$ 0.00	11.40 $\pm$ 0.00	10.20 $\pm$ 0.00	10.60 $\pm$ 0.00	11.20 $\pm$ 0.00	10.92 $\pm$ 0.00	10.92 $\pm$ 0.00
antiochoid	62.55 $\pm$ 0.04	61.80 $\pm$ 0.00	60.67 $\pm$ 0.00	58.99 $\pm$ 0.00	58.80 $\pm$ 0.00	60.40 $\pm$ 0.00	61.99 $\pm$ 0.00	60.40 $\pm$ 0.00	58.43 $\pm$ 0.00	58.24 $\pm$ 0.00	56.74 $\pm$ 0.00	59.18 $\pm$ 0.00
backdoor	64.15 $\pm$ 0.04	52.30 $\pm$ 1.87	40.62 $\pm$ 1.46	30.25 $\pm$ 1.34	20.98 $\pm$ 1.20	42.86 $\pm$ 1.32	52.53 $\pm$ 1.63	40.37 $\pm$ 2.04	28.71 $\pm$ 1.50	20.21 $\pm$ 0.78	17.52 $\pm$ 0.83	31.87 $\pm$ 1.22
breastw	96.72 $\pm$ 0.64	96.23 $\pm$ 0.39	96.17 $\pm$ 0.47	<u>95.99</u> $\pm$ 0.39	95.99 $\pm$ 0.39	96.22 $\pm$ 0.43	96.00 $\pm$ 0.44	96.05 $\pm$ 0.33	<u>95.87</u> $\pm$ 0.28	95.99 $\pm$ 0.39	95.93 $\pm$ 0.32	95.97 $\pm$ 0.31
campaign	99.94 $\pm$ 0.00	50.62 $\pm$ 0.00	51.14 $\pm$ 0.00	51.38 $\pm$ 0.00	51.37 $\pm$ 0.00	50.33 $\pm$ 0.00	50.37 $\pm$ 0.00	50.80 $\pm$ 0.00	51.70 $\pm$ 0.00	51.29 $\pm$ 0.00	51.10 $\pm$ 0.00	51.10 $\pm$ 0.00
cardio	63.64 $\pm$ 0.00	61.26 $\pm$ 0.00	61.98 $\pm$ 0.00	63.64 $\pm$ 0.00	64.20 $\pm$ 0.00	62.95 $\pm$ 0.00	60.93 $\pm$ 0.00	60.83 $\pm$ 0.00	64.20 $\pm$ 0.00	67.61 $\pm$ 0.00	69.32 $\pm$ 0.00	65.00 $\pm$ 0.00
cardiologocography	44.64 $\pm$ 0.00	45.71 $\pm$ 0.00	47.00 $\pm$ 0.00	48.50 $\pm$ 0.00	49.14 $\pm$ 0.00	47.00 $\pm$ 0.00	46.33 $\pm$ 0.00	46.78 $\pm$ 0.00	47.85 $\pm$ 0.00	50.86 $\pm$ 0.00	51.93 $\pm$ 0.00	48.76 $\pm$ 0.00
celeba	15.83 $\pm$ 0.69	17.05 $\pm$ 0.43	18.17 $\pm$ 0.61	19.02 $\pm$ 0.69	19.30 $\pm$ 0.60	17.87 $\pm$ 0.57	<u>17.08</u> $\pm$ 0.58	18.41 $\pm$ 0.65	19.30 $\pm$ 0.61	20.27 $\pm$ 0.68	20.48 $\pm$ 0.68	19.11 $\pm$ 0.61
census	22.22 $\pm$ 0.54	21.93 $\pm$ 0.52	21.46 $\pm$ 0.25	21.38 $\pm$ 0.48	21.12 $\pm$ 0.20	21.62 $\pm$ 0.14	22.23 $\pm$ 0.42	21.48 $\pm$ 0.40	21.47 $\pm$ 0.57	21.33 $\pm$ 0.65	21.26 $\pm$ 0.50	21.55 $\pm$ 0.24
cover	69.15 $\pm$ 2.12	66.87 $\pm$ 2.35	63.15 $\pm$ 2.07	55.99 $\pm$ 2.08	53.06 $\pm$ 2.23	61.65 $\pm$ 2.14	65.04 $\pm$ 1.92	60.56 $\pm$ 2.04	52.76 $\pm$ 1.83	42.69 $\pm$ 1.94	39.02 $\pm$ 1.99	52.19 $\pm$ 1.92
donors	97.27 $\pm$ 0.36	96.20 $\pm$ 0.45	94.49 $\pm$ 0.55	91.70 $\pm$ 0.90	90.57 $\pm$ 0.86	94.05 $\pm$ 0.60	94.98 $\pm$ 0.62	92.36 $\pm$ 0.59	88.71 $\pm$ 0.99	80.36 $\pm$ 1.90	76.02 $\pm$ 1.50	86.60 $\pm$ 0.98
fault	56.02 $\pm$ 0.00	55.72 $\pm$ 0.00	55.57 $\pm$ 0.00	56.91 $\pm$ 0.00	57.36 $\pm$ 0.00	56.32 $\pm$ 0.00	55.57 $\pm$ 0.00	55.87 $\pm$ 0.00	57.06 $\pm$ 0.00	57.50 $\pm$ 0.00	58.25 $\pm$ 0.00	56.85 $\pm$ 0.00
fraud	48.18 $\pm$ 4.56	49.60 $\pm$ 3.28	49.00 $\pm$ 3.95	46.66 $\pm$ 3.52	45.46 $\pm$ 3.60	47.78 $\pm$ 3.53	47.78 $\pm$ 3.09	49.39 $\pm$ 5.24	46.64 $\pm$ 4.18	42.58 $\pm$ 3.16	41.64 $\pm$ 3.36	45.61 $\pm$ 3.48
glass	47.81 $\pm$ 5.78	35.14 $\pm$ 2.75	27.98 $\pm$ 2.40	18.37 $\pm$ 2.40	17.81 $\pm$ 2.98	29.42 $\pm$ 2.61	29.87 $\pm$ 9.62	22.55 $\pm$ 6.87	18.58 $\pm$ 4.19	17.23 $\pm$ 3.73	17.23 $\pm$ 3.73	21.09 $\pm$ 5.16
hepatitis	98.94 $\pm$ 1.41	94.16 $\pm$ 2.13	81.68 $\pm$ 4.18	75.97 $\pm$ 4.78	71.90 $\pm$ 4.14	84.54 $\pm$ 2.23	81.98 $\pm$ 4.50	60.39 $\pm$ 6.21	62.31 $\pm$ 6.09	60.39 $\pm$ 6.21	60.04 $\pm$ 7.30	66.15 $\pm$ 4.71
lftp	98.50 $\pm$ 0.38	95.10 $\pm$ 2.50	88.26 $\pm$ 1.38	82.85 $\pm$ 3.80	78.96 $\pm$ 6.40	88.73 $\pm$ 2.54	100.00 $\pm$ 0.00	58.57 $\pm$ 0.86	92.67 $\pm$ 0.91	92.67 $\pm$ 0.91	92.67 $\pm$ 0.91	95.39 $\pm$ 0.75
indb	5.20 $\pm$ 0.00	5.40 $\pm$ 0.00	5.40 $\pm$ 0.00	5.20 $\pm$ 0.00	5.20 $\pm$ 0.00	5.20 $\pm$ 0.00	5.40 $\pm$ 0.00	5.40 $\pm$ 0.00	5.40 $\pm$ 0.00	5.00 $\pm$ 0.00	5.00 $\pm$ 0.00	5.24 $\pm$ 0.00
intercards	55.16 $\pm$ 0.00	51.63 $\pm$ 0.00	48.37 $\pm$ 0.00	46.47 $\pm$ 0.00	46.20 $\pm$ 0.00	49.12 $\pm$ 0.00	51.90 $\pm$ 0.00	45.11 $\pm$ 0.00	45.11 $\pm$ 0.00	45.11 $\pm$ 0.00	46.68 $\pm$ 0.00	46.68 $\pm$ 0.00
ionosphere	92.33 $\pm$ 1.17	92.05 $\pm$ 1.63	91.63 $\pm$ 1.09	90.41 $\pm$ 1.35	89.19 $\pm$ 1.72	91.57 $\pm$ 1.24	91.23 $\pm$ 1.86	91.81 $\pm$ 1.87	89.15 $\pm$ 1.91	85.06 $\pm$ 3.37	82.75 $\pm$ 3.12	87.86 $\pm$ 1.86
letter	52.29 $\pm$ 0.00	51.24 $\pm$ 0.00	49.06 $\pm$ 0.00	45.99 $\pm$ 0.00	45.99 $\pm$ 0.00	48.70 $\pm$ 0.00	51.46 $\pm$ 0.00	49.20 $\pm$ 0.00	45.76 $\pm$ 0.00	42.69 $\pm$ 0.00	41.34 $\pm$ 0.00	46.11 $\pm$ 0.00
letter	1.00 $\pm$ 0.00	1.00 $\pm$ 0.00	1.00 $\pm$ 0.00	1.00 $\pm$ 0.00	1.00 $\pm$ 0.00	1.00 $\pm$ 0.00	1.00 $\pm$ 0.00	1.00 $\pm$ 0.00	1.00 $\pm$ 0.00	1.00 $\pm$ 0.00	1.00 $\pm$ 0.00	1.00 $\pm$ 0.00
lymphography	97.89 $\pm$ 2.21	93.61 $\pm$ 7.97	94.67 $\pm$ 6.53	92.71 $\pm$ 6.09	91.68 $\pm$ 7.34	94.11 $\pm$ 5.93	91.57 $\pm$ 5.96	90.69 $\pm$ 3.65	90.69 $\pm$ 3.65	92.96 $\pm$ 1.97	92.96 $\pm$ 1.97	91.77 $\pm$ 3.69
magic gamma	76.79 $\pm$ 0.00	76.20 $\pm$ 0.00	75.49 $\pm$ 0.00	74.84 $\pm$ 0.00	74.17 $\pm$ 0.00	75.61 $\pm$ 0.00	76.17 $\pm$ 0.00	75.13 $\pm$ 0.00	74.89 $\pm$ 0.00	73.00 $\pm$ 0.00	73.00 $\pm$ 0.00	74.48 $\pm$ 0.00
maninography	41.02 $\pm$ 0.00	41.15 $\pm$ 0.00	44.23 $\pm$ 0.00	44.23 $\pm$ 0.00	44.23 $\pm$ 0.00	43.35 $\pm$ 0.00	40.38 $\pm$ 0.00	43.46 $\pm$ 0.00	43.85 $\pm$ 0.00	43.85 $\pm$ 0.00	43.46 $\pm$ 0.00	43.00 $\pm$ 0.00
mnist	72.71 $\pm$ 0.00	72.29 $\pm$ 0.00	71.57 $\pm$ 0.00	70.43 $\pm$ 0.00	69.86 $\pm$ 0.00	71.37 $\pm$ 0.00	71.86 $\pm$ 0.00	71.29 $\pm$ 0.00	100.00 $\pm$ 0.00	69.71 $\pm$ 0.00	69.71 $\pm$ 0.00	70.49 $\pm$ 0.00
mnist	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00
mnist	30.00 $\pm$ 0.00	24.00 $\pm$ 0.00	12.00 $\pm$ 0.00	7.33 $\pm$ 0.00	6.07 $\pm$ 0.00	16.00 $\pm$ 0.00	31.33 $\pm$ 0.00	39.22 $\pm$ 0.00	60.20 $\pm$ 0.00	30.08 $\pm$ 0.00	26.27 $\pm$ 0.00	8.35 $\pm$ 0.00
mnist	94.23 $\pm$ 0.00	92.31 $\pm$ 0.00	91.03 $\pm$ 0.00	80.13 $\pm$ 0.00	78.18 $\pm$ 0.00	87.18 $\pm$ 0.00	90.38 $\pm$ 0.00	90.38 $\pm$ 0.00	73.72 $\pm$ 0.00	61.74 $\pm$ 0.00	62.82 $\pm$ 0.00	58.16 $\pm$ 0.00
mnist	74.73 $\pm$ 2.13	72.94 $\pm$ 2.46	71.48 $\pm$ 1.96	71.41 $\pm$ 2.36	71.03 $\pm$ 1.99	72.30 $\pm$ 1.97	71.03 $\pm$ 2.53	70.18 $\pm$ 2.12	71.03 $\pm$ 2.12	71.67 $\pm$ 2.10	72.03 $\pm$ 2.02	71.30 $\pm$ 2.00
mnist	72.20 $\pm$ 0.00	71.76 $\pm$ 0.00	70.68 $\pm$ 0.00	69.79 $\pm$ 0.00	69.20 $\pm$ 0.00	70.73 $\pm$ 0.00	71.81 $\pm$ 0.00	91.53 $\pm$ 0.00	69.84 $\pm$ 0.00	67.93 $\pm$ 0.00	67.29 $\pm$ 0.00	69.52 $\pm$ 0.00
mnist	90.14 $\pm$ 0.00	90.14 $\pm$ 0.00	90.14 $\pm$ 0.00	92.96 $\pm$ 0.00	92.96 $\pm$ 0.00	91.27 $\pm$ 0.00	90.11 $\pm$ 0.00	91.53 $\pm$ 0.00	92.96 $\pm$ 0.00	92.96 $\pm$ 0.00	92.96 $\pm$ 0.00	92.11 $\pm$ 0.00
mnist	98.35 $\pm$ 0.00	98.23 $\pm$ 0.00	98.15 $\pm$ 0.00	98.12 $\pm$ 0.00	98.12 $\pm$ 0.00	98.19 $\pm$ 0.00	98.23 $\pm$ 0.00	98.05 $\pm$ 0.00	98.05 $\pm$ 0.00	98.09 $\pm$ 0.00	98.11 $\pm$ 0.00	98.11 $\pm$ 0.00
mnist	97.12 $\pm$ 0.46	96.83 $\pm$ 0.38	96.14 $\pm$ 0.17	94.73 $\pm$ 0.23	94.30 $\pm$ 0.27	95.82 $\pm$ 0.14	96.71 $\pm$ 0.41	95.85 $\pm$ 0.17	94.56 $\pm$ 0.29	91.73 $\pm$ 0.16	90.00 $\pm$ 0.28	93.89 $\pm$ 0.13
mnist	68.05 $\pm$ 5.12	68.05 $\pm$ 5.12	69.59 $\pm$ 3.95	69.59 $\pm$ 3.95	69.33 $\pm$ 8.31	67.73 $\pm$ 4.99	69.59 $\pm$ 3.95	69.59 $\pm$ 3.95	69.59 $\pm$ 3.95	69.59 $\pm$ 3.95	69.59 $\pm$ 3.95	69.59 $\pm$ 3.95
mnist	80.88 $\pm$ 0.00	80.29 $\pm$ 0.00	80.23 $\pm$ 0.00	79.81 $\pm$ 0.00	79.45 $\pm$ 0.00	80.13 $\pm$ 0.00	80.52 $\pm$ 0.00	79.93 $\pm$ 0.00	79.15 $\pm$ 0.00	78.98 $\pm$ 0.00	78.80 $\pm$ 0.00	79.48 $\pm$ 0.00
mnist	3.28 $\pm$ 0.00	3.28 $\pm$ 0.00	1.61 $\pm$ 0.00	1.64 $\pm$ 0.00	2.62 $\pm$ 0.00	3.28 $\pm$ 0.00	3.28 $\pm$ 0.00	3.28 $\pm$ 0.00	3.28 $\pm$ 0.00	3.28 $\pm$ 0.00	3.28 $\pm$ 0.00	3.28 $\pm$ 0.00
mnist	85.93 $\pm$ 2.66	80.63 $\pm$ 2.92	74.04 $\pm$ 4.45	68.12 $\pm$ 7.14	66.98 $\pm$ 6.79	75.14 $\pm$ 4.45	75.57 $\pm$ 8.41	68.89 $\pm$ 8.23	62.67 $\pm$ 6.19	60.52 $\pm$ 7.08	61.78 $\pm$ 6.58	66.49 $\pm$ 7.15
mnist	75.27 $\pm$ 0.00	75.27 $\pm$ 0.00	74.19 $\pm$ 0.00	74.19 $\pm$ 0.00	74.19 $\pm$ 0.00	74.62 $\pm$ 0.00	75.27 $\pm$ 0.00	73.12 $\pm$ 0.00	72.04 $\pm$ 0.00	73.12 $\pm$ 0.00	74.19 $\pm$ 0.00	73.55 $\pm$ 0.00
mnist	49.03 $\pm$ 4.93	32.73 $\pm$ 5.11	24.06 $\pm$ 3.90	15.07 $\pm$ 1.64	12.31 $\pm$ 2.44	26.68 $\pm$ 3.24	23.02 $\pm$ 5.17	17.18 $\pm$ 2.01	12.19 $\pm$ 3.52	9.07 $\pm$ 3.30	8.51 $\pm$ 2.65	13.99 $\pm$ 2.95
mnist	28.00 $\pm$ 0.00	28.00 $\pm$ 0.00	28.00 $\pm$ 0.00	30.00 $\pm$ 0.00	30.00 $\pm$ 0.00	28.00 $\pm$ 0.00	30.00 $\pm$ 0.00	26.00 $\pm$ 0.00	26.00 $\pm$ 0.00	26.00 $\pm$ 0.00	26.00 $\pm$ 0.00	26.00 $\pm$ 0.00
mnist	26.00 $\pm$ 0.00	26.00 $\pm$ 0.00	27.00 $\pm$ 0.00	28.00 $\pm$ 0.00	28.00 $\pm$ 0.00	27.00 $\pm$ 0.00	27.00 $\pm$ 0.00	27.00 $\pm$ 0.00	27.00 $\pm$ 0.00	27.00 $\pm$ 0.00	27.00 $\pm$ 0.00	27.00 $\pm$ 0.00
mnist	89.35 $\pm$ 3.08	88.05 $\pm$ 3.35	88.05 $\pm$ 3.35	89.16 $\pm$ 2.38	89.16 $\pm$ 2.38	88.75 $\pm$ 2.48	83.65 $\pm$ 4.53	84.82 $\pm$ 3.86	83.72 $\pm$ 5.17	89.16 $\pm$ 2.38	89.16 $\pm$ 2.38	86.10 $\pm$ 2.43
mnist	85.05 $\pm$ 6.11	83.36 $\pm$ 7.16	81.00 $\pm$ 4.08	81.00 $\pm$ 4.08	81.00 $\pm$ 4.08	82.28 $\pm$ 4.91	80.92 $\pm$ 4.47	78.72 $\pm$ 5.61	79.49 $\pm$ 5.02	77.95 $\pm$ 6.52	76.84 $\pm$ 4.81</	

Table 16.2: Average F1 score $\pm$ standard dev. over five seeds for the semi-supervised setting of ICL and DTE-C baselines with varying hyperparameter (HP) values; For ICL, the learning rate $\in \{0.1, 0.02, 0.001, 0.0001, 1e-05\}$, for DTE-C, $k \in \{5, 10, 20, 40, 50\}$. Also reported is the avg model. We use **bold** and underline respectively to mark the **best** and the **worst** performance of each model to showcase the variability of performance across different HP settings.

dataset	ICL-0.1	ICL-0.01	ICL-0.001	ICL-0.0001	ICL-1e-05	ICL-avg	DTE-C-5	DTE-C-10	DTE-C-20	DTE-C-40	DTE-C-50	DTE-C-avg
abi	4.51 $\pm$ 0.69	4.34 $\pm$ 0.42	5.28 $\pm$ 0.47	4.68 $\pm$ 0.30	4.16 $\pm$ 0.38	4.59 $\pm$ 0.07	4.75 $\pm$ 0.27	4.27 $\pm$ 0.19	4.28 $\pm$ 0.10	4.51 $\pm$ 0.17	0.00 $\pm$ 0.00	3.56 $\pm$ 0.03
amazon	10.44 $\pm$ 0.46	9.76 $\pm$ 0.34	9.92 $\pm$ 0.84	10.08 $\pm$ 0.35	9.52 $\pm$ 0.43	9.94 $\pm$ 0.32	11.96 $\pm$ 0.37	11.96 $\pm$ 0.37	11.60 $\pm$ 0.27	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	7.01 $\pm$ 0.72
amthroid	54.87 $\pm$ 13.24	42.25 $\pm$ 3.55	53.45 $\pm$ 4.13	54.72 $\pm$ 5.45	57.53 $\pm$ 2.97	52.56 $\pm$ 3.29	77.94 $\pm$ 0.28	77.94 $\pm$ 0.28	77.53 $\pm$ 0.85	75.43 $\pm$ 0.93	0.00 $\pm$ 0.00	61.68 $\pm$ 0.20
backdoor	87.17 $\pm$ 0.98	87.32 $\pm$ 0.99	87.11 $\pm$ 0.09	86.85 $\pm$ 0.95	85.37 $\pm$ 1.01	86.76 $\pm$ 1.01	83.03 $\pm$ 1.01	84.50 $\pm$ 0.60	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	42.75 $\pm$ 1.54
breastw	95.98 $\pm$ 0.34	96.07 $\pm$ 0.94	96.80 $\pm$ 0.40	97.44 $\pm$ 0.55	96.48 $\pm$ 0.28	96.48 $\pm$ 0.28	88.50 $\pm$ 1.59	90.10 $\pm$ 1.35	92.46 $\pm$ 1.78	95.31 $\pm$ 0.70	96.11 $\pm$ 0.44	92.50 $\pm$ 0.79
campaign	48.12 $\pm$ 0.36	46.81 $\pm$ 1.72	50.68 $\pm$ 0.66	51.37 $\pm$ 0.85	53.40 $\pm$ 0.55	50.07 $\pm$ 0.49	51.98 $\pm$ 0.70	52.45 $\pm$ 1.07	52.33 $\pm$ 1.00	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	31.35 $\pm$ 0.44
cardio	49.09 $\pm$ 11.28	61.93 $\pm$ 5.57	58.86 $\pm$ 1.59	57.95 $\pm$ 2.80	40.57 $\pm$ 4.28	53.68 $\pm$ 2.83	58.30 $\pm$ 0.58	57.84 $\pm$ 0.43	58.07 $\pm$ 0.23	0.34 $\pm$ 0.68	0.00 $\pm$ 0.00	34.91 $\pm$ 0.09
cardiolography	36.14 $\pm$ 1.28	41.07 $\pm$ 1.73	39.18 $\pm$ 4.49	35.36 $\pm$ 2.14	32.66 $\pm$ 1.59	36.88 $\pm$ 1.27	39.91 $\pm$ 1.05	39.48 $\pm$ 1.54	37.73 $\pm$ 1.73	62.02 $\pm$ 0.00	62.02 $\pm$ 0.00	48.28 $\pm$ 0.39
celeba	15.42 $\pm$ 2.29	17.97 $\pm$ 2.55	17.20 $\pm$ 1.92	17.46 $\pm$ 1.17	16.17 $\pm$ 0.65	16.84 $\pm$ 0.50	19.18 $\pm$ 2.74	17.12 $\pm$ 1.45	14.31 $\pm$ 2.28	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	10.12 $\pm$ 1.04
cover	22.72 $\pm$ 1.73	24.06 $\pm$ 2.05	25.80 $\pm$ 1.34	27.15 $\pm$ 1.12	24.06 $\pm$ 1.31	24.76 $\pm$ 0.50	17.58 $\pm$ 1.42	17.54 $\pm$ 1.79	16.44 $\pm$ 1.41	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	10.31 $\pm$ 0.52
cover	26.77 $\pm$ 15.24	15.53 $\pm$ 8.17	42.68 $\pm$ 17.00	53.70 $\pm$ 16.88	44.34 $\pm$ 20.24	36.61 $\pm$ 2.59	76.51 $\pm$ 2.37	68.92 $\pm$ 4.22	46.54 $\pm$ 3.93	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	38.39 $\pm$ 0.62
donors	60.33 $\pm$ 3.36	81.85 $\pm$ 3.56	83.59 $\pm$ 4.47	89.28 $\pm$ 2.66	92.77 $\pm$ 1.39	75.24 $\pm$ 1.61	87.99 $\pm$ 1.87	75.05 $\pm$ 7.18	63.08 $\pm$ 1.70	11.80 $\pm$ 23.60	0.00 $\pm$ 0.00	47.58 $\pm$ 4.61
fault	57.54 $\pm$ 10.13	59.52 $\pm$ 2.09	58.57 $\pm$ 1.11	58.37 $\pm$ 0.79	59.13 $\pm$ 1.36	58.87 $\pm$ 1.51	55.57 $\pm$ 1.57	55.16 $\pm$ 0.81	55.33 $\pm$ 1.66	97.03 $\pm$ 0.40	96.91 $\pm$ 0.24	72.00 $\pm$ 0.47
fraud	43.53 $\pm$ 20.21	48.05 $\pm$ 8.02	58.90 $\pm$ 6.77	66.88 $\pm$ 4.88	79.18 $\pm$ 3.21	62.39 $\pm$ 4.91	75.61 $\pm$ 1.76	74.25 $\pm$ 0.50	22.55 $\pm$ 24.73	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	30.48 $\pm$ 4.66
glass	57.53 $\pm$ 20.21	57.05 $\pm$ 16.03	84.05 $\pm$ 6.11	87.24 $\pm$ 5.04	82.30 $\pm$ 5.41	70.87 $\pm$ 3.12	93.43 $\pm$ 5.07	34.48 $\pm$ 4.93	31.36 $\pm$ 0.69	19.45 $\pm$ 5.66	0.00 $\pm$ 0.00	24.15 $\pm$ 2.73
hepatitis	99.64 $\pm$ 0.71	94.69 $\pm$ 7.81	90.64 $\pm$ 0.71	90.64 $\pm$ 0.71	99.64 $\pm$ 0.71	98.65 $\pm$ 2.11	96.40 $\pm$ 3.01	94.63 $\pm$ 4.00	92.51 $\pm$ 3.12	51.68 $\pm$ 0.78	18.97 $\pm$ 10.60	70.84 $\pm$ 3.40
htrp	93.91 $\pm$ 10.35	96.07 $\pm$ 3.07	97.69 $\pm$ 1.51	99.36 $\pm$ 0.19	97.28 $\pm$ 0.15	98.08 $\pm$ 1.97	98.08 $\pm$ 1.97	98.08 $\pm$ 1.97	18.33 $\pm$ 10.68	16.75 $\pm$ 12.06	0.00 $\pm$ 0.00	24.79 $\pm$ 12.88
imdb	10.52 $\pm$ 0.65	10.44 $\pm$ 0.54	9.84 $\pm$ 0.37	9.76 $\pm$ 0.15	10.19 $\pm$ 0.23	6.64 $\pm$ 0.89	8.40 $\pm$ 1.23	7.39 $\pm$ 1.04	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	4.47 $\pm$ 0.43
internetids	55.92 $\pm$ 2.66	57.45 $\pm$ 0.61	57.77 $\pm$ 1.10	58.26 $\pm$ 0.50	49.02 $\pm$ 1.45	55.68 $\pm$ 0.94	67.99 $\pm$ 1.98	65.87 $\pm$ 0.95	64.95 $\pm$ 2.24	41.85 $\pm$ 0.00	41.85 $\pm$ 0.00	56.50 $\pm$ 0.87
ionosphere	92.64 $\pm$ 4.66	91.43 $\pm$ 4.67	93.86 $\pm$ 1.63	94.48 $\pm$ 0.71	94.49 $\pm$ 1.56	93.38 $\pm$ 2.21	89.67 $\pm$ 1.44	89.41 $\pm$ 1.53	80.72 $\pm$ 1.37	78.12 $\pm$ 2.07	49.44 $\pm$ 16.82	79.23 $\pm$ 1.14
landsat	49.50 $\pm$ 1.24	47.94 $\pm$ 1.88	54.51 $\pm$ 0.52	51.25 $\pm$ 0.71	47.97 $\pm$ 1.09	50.83 $\pm$ 0.68	35.45 $\pm$ 0.70	35.47 $\pm$ 0.73	31.72 $\pm$ 3.08	51.29 $\pm$ 4.34	54.46 $\pm$ 0.00	41.47 $\pm$ 2.39
letter	6.80 $\pm$ 4.21	4.00 $\pm$ 1.55	3.60 $\pm$ 1.26	3.20 $\pm$ 0.75	11.60 $\pm$ 2.06	10.00 $\pm$ 0.00	2.49 $\pm$ 1.50	3.00 $\pm$ 1.55	3.20 $\pm$ 1.57	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	1.72 $\pm$ 0.37
lymphography	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	77.11 $\pm$ 6.79	70.84 $\pm$ 1.53	71.09 $\pm$ 2.69	60.09 $\pm$ 2.69	0.00 $\pm$ 0.00	58.34 $\pm$ 6.75
magic-gamma	61.88 $\pm$ 2.50	69.92 $\pm$ 2.08	69.99 $\pm$ 1.87	70.31 $\pm$ 0.50	71.64 $\pm$ 0.63	69.34 $\pm$ 1.08	70.82 $\pm$ 0.48	80.19 $\pm$ 0.73	80.73 $\pm$ 1.81	89.73 $\pm$ 1.81	96.10 $\pm$ 0.00	88.36 $\pm$ 1.69
maninography	36.02 $\pm$ 8.76	38.15 $\pm$ 8.41	27.69 $\pm$ 1.80	23.98 $\pm$ 2.63	18.15 $\pm$ 1.23	29.34 $\pm$ 2.59	32.69 $\pm$ 5.73	34.62 $\pm$ 2.19	35.08 $\pm$ 2.68	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	38.42 $\pm$ 1.48
mnist	35.57 $\pm$ 1.84	46.20 $\pm$ 3.18	50.54 $\pm$ 1.26	59.20 $\pm$ 1.30	62.66 $\pm$ 2.09	52.78 $\pm$ 0.77	63.80 $\pm$ 1.77	53.74 $\pm$ 3.75	53.74 $\pm$ 3.75	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	61.11 $\pm$ 5.2
mnist	29.57 $\pm$ 5.08	10.00 $\pm$ 0.00	100.00 $\pm$ 0.00	100.00 $\pm$ 0.00	46.00 $\pm$ 5.11	89.36 $\pm$ 1.65	14.80 $\pm$ 1.81	6.40 $\pm$ 3.88	9.33 $\pm$ 5.72	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	60.00 $\pm$ 0.00
optdigits	52.10 $\pm$ 2.81	64.08 $\pm$ 2.96	62.31 $\pm$ 2.65	66.03 $\pm$ 5.18	62.20 $\pm$ 1.05	62.38 $\pm$ 0.29	63.97 $\pm$ 0.36	51.36 $\pm$ 7.40	43.46 $\pm$ 7.50	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	48.38 $\pm$ 0.37
pagoblocks	60.51 $\pm$ 10.58	60.51 $\pm$ 10.58	59.29 $\pm$ 5.15	63.88 $\pm$ 1.08	51.03 $\pm$ 3.57	56.50 $\pm$ 3.52	63.97 $\pm$ 0.36	51.36 $\pm$ 7.40	43.46 $\pm$ 7.50	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	32.96 $\pm$ 1.43
plma	68.77 $\pm$ 4.96	68.22 $\pm$ 2.03	74.95 $\pm$ 1.27	71.40 $\pm$ 1.67	68.83 $\pm$ 2.08	73.62 $\pm$ 2.09	72.06 $\pm$ 0.60	72.97 $\pm$ 0.77	64.71 $\pm$ 3.39	67.99 $\pm$ 2.34	63.96 $\pm$ 1.44	65.74 $\pm$ 2.10
satellite	65.94 $\pm$ 10.44	72.05 $\pm$ 2.96	76.27 $\pm$ 0.81	78.70 $\pm$ 0.90	74.22 $\pm$ 0.97	69.69 $\pm$ 7.51	76.31 $\pm$ 3.03	60.85 $\pm$ 5.87	46.20 $\pm$ 5.07	91.51 $\pm$ 9.02	96.02 $\pm$ 0.00	81.04 $\pm$ 1.96
satimage-2	38.03 $\pm$ 8.39	72.08 $\pm$ 3.26	91.83 $\pm$ 0.56	89.58 $\pm$ 1.13	56.34 $\pm$ 12.38	98.38 $\pm$ 0.14	98.91 $\pm$ 0.12	97.98 $\pm$ 0.00	97.73 $\pm$ 0.10	92.83 $\pm$ 2.84	0.00 $\pm$ 0.00	42.31 $\pm$ 7.14
shuttle	97.17 $\pm$ 1.11	98.27 $\pm$ 0.10	98.83 $\pm$ 0.14	98.91 $\pm$ 0.12	98.39 $\pm$ 0.11	96.43 $\pm$ 0.55	82.15 $\pm$ 0.39	82.23 $\pm$ 0.37	81.06 $\pm$ 0.57	80.48 $\pm$ 0.35	54.71 $\pm$ 0.80	77.31 $\pm$ 0.58
skin	39.28 $\pm$ 28.70	59.49 $\pm$ 7.84	54.99 $\pm$ 13.50	48.81 $\pm$ 9.06	50.03 $\pm$ 10.03	77.87 $\pm$ 0.46	69.59 $\pm$ 3.35	69.59 $\pm$ 3.35	49.07 $\pm$ 5.05	27.32 $\pm$ 14.61	0.00 $\pm$ 0.00	76.25 $\pm$ 0.32
smtp	76.97 $\pm$ 2.12	76.88 $\pm$ 3.76	80.35 $\pm$ 0.48	80.23 $\pm$ 0.62	74.93 $\pm$ 0.44	77.87 $\pm$ 0.46	80.19 $\pm$ 0.18	80.27 $\pm$ 0.41	80.56 $\pm$ 0.29	87.61 $\pm$ 0.00	87.61 $\pm$ 0.00	83.25 $\pm$ 0.07
spambase	2.95 $\pm$ 1.23	3.61 $\pm$ 1.23	3.28 $\pm$ 1.04	2.95 $\pm$ 1.21	4.59 $\pm$ 1.61	3.48 $\pm$ 0.92	3.93 $\pm$ 0.80	2.95 $\pm$ 1.61	3.28 $\pm$ 1.80	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	2.03 $\pm$ 0.25
speech	31.81 $\pm$ 12.18	42.45 $\pm$ 13.97	83.03 $\pm$ 3.34	76.24 $\pm$ 6.92	73.47 $\pm$ 7.50	62.00 $\pm$ 6.23	63.62 $\pm$ 0.79	57.63 $\pm$ 9.24	56.93 $\pm$ 11.83	75.93 $\pm$ 10.89	50.23 $\pm$ 10.89	60.69 $\pm$ 0.57
stamps	68.39 $\pm$ 4.17	61.29 $\pm$ 1.80	61.08 $\pm$ 3.63	63.87 $\pm$ 3.57	33.76 $\pm$ 5.95	57.68 $\pm$ 1.74	73.76 $\pm$ 2.21	76.13 $\pm$ 1.72	75.27 $\pm$ 0.00	78.28 $\pm$ 0.80	0.00 $\pm$ 0.00	66.69 $\pm$ 0.57
thyroid	23.17 $\pm$ 6.96	29.15 $\pm$ 9.05	52.53 $\pm$ 11.60	64.81 $\pm$ 3.19	54.89 $\pm$ 6.76	44.91 $\pm$ 2.82	43.88 $\pm$ 5.22	36.26 $\pm$ 1.73	32.59 $\pm$ 10.94	21.68 $\pm$ 8.71	3.48 $\pm$ 6.96	28.18 $\pm$ 7.79
vertebral	22.40 $\pm$ 17.59	18.80 $\pm$ 6.52	30.80 $\pm$ 7.44	30.80 $\pm$ 12.43	24.00 $\pm$ 7.04	25.36 $\pm$ 2.94	40.00 $\pm$ 5.51	40.00 $\pm$ 5.51	45.20 $\pm$ 0.98	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	25.20 $\pm$ 1.36
vowels	38.20 $\pm$ 17.57	35.40 $\pm$ 7.47	11.69 $\pm$ 1.62	28.40 $\pm$ 4.22	32.28 $\pm$ 5.79	28.40 $\pm$ 4.22	11.20 $\pm$ 1.72	11.80 $\pm$ 2.14	12.20 $\pm$ 1.60	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	7.04 $\pm$ 0.23
waveform	47.80 $\pm$ 3.06	84.03 $\pm$ 7.23	96.89 $\pm$ 4.35	95.71 $\pm$ 4.90	90.32 $\pm$ 5.34	89.16 $\pm$ 3.80	40.59 $\pm$ 7.11	34.81 $\pm$ 6.65	40.06 $\pm$ 12.93	67.85 $\pm$ 5.79	0.00 $\pm$ 0.00	36.84 $\pm$ 4.92
wbc	78.89 $\pm$ 7.89	80.42 $\pm$ 5.68	84.10 $\pm$ 6.72	89.61 $\pm$ 4.50	78.33 $\pm$ 5.25	78.86 $\pm$ 2.34	78.62 $\pm$ 5.69	70.31 $\pm$ 4.12	5.37 $\pm$ 2.97	16.96 $\pm$ 5.74	0.00 $\pm$ 0.00	47.39 $\pm$ 3.91
wdbc	64.32 $\pm$ 18.21	11.36 $\pm$ 4.03	37.04 $\pm$ 8.60	81.10 $\pm$ 3.50	37.59 $\pm$ 2.52	26.35 $\pm$ 2.35	19.53 $\pm$ 2.50	95.05 $\pm$ 5.74	97.27 $\pm$ 3.64	41.73 $\pm$ 22.73		

Table 18.1: Average AUPR $\pm$ standard dev. over five seeds for the semi-supervised setting on ADBench. Rank of each model per dataset is provided (in parentheses) (the lower, the better). We use blue and green respectively to mark the top-1 and the top-2 method

[illegible]

Table 18.2: Average AUPR $\pm$ standard dev. over five seeds for the semi-supervised setting on ADBench. Rank of each model per dataset is provided (in parentheses) (the lower, the better). We use blue and green respectively to mark the top-1 and the top-2 method.

[illegible]

Table 19.1: Average F1 score $\pm$ standard dev. over five seeds for the semi-supervised setting on ADBench. Rank of each model per dataset is provided (in parentheses) (the lower, the better). We use blue and green respectively to mark the top-1 and the top-2 method.

[illegible]

J Benchmark OD Datasets

Table 20: Description of all datasets in ADBench Livernoche et al. (2024). Datasets in [blue](#) are image and text datasets that are vectorized through pretrained encoders. We refer to the original paper for details.

Dataset Name	# Samples	# Features	# Anomaly	% Anomaly	Category
ALOI	49534	27	1508	3.04	Image
annthyroid	7200	6	534	7.42	Healthcare
backdoor	95329	196	2329	2.44	Network
breastw	683	9	239	34.99	Healthcare
campaign	41188	62	4640	11.27	Finance
cardio	1831	21	176	9.61	Healthcare
Cardiotocography	2114	21	466	22.04	Healthcare
celeba	202599	39	4547	2.24	Image
census	299285	500	18568	6.20	Sociology
cover	286048	10	2747	0.96	Botany
donors	619326	10	36710	5.93	Sociology
fault	1941	27	673	34.67	Physical
fraud	284807	29	492	0.17	Finance
glass	214	7	9	4.21	Forensic
Hepatitis	80	19	13	16.25	Healthcare
http	567498	3	2211	0.39	Web
InternetAds	1966	1555	368	18.72	Image
Ionosphere	351	32	126	35.90	Oryctognosy
landsat	6435	36	1333	20.71	Astronautics
letter	1600	32	100	6.25	Image
Lymphography	148	18	6	4.05	Healthcare
magic.gamma	19020	10	6688	35.16	Physical
mammography	11183	6	260	2.32	Healthcare
mnist	7603	100	700	9.21	Image
musk	3062	166	97	3.17	Chemistry
optdigits	5216	64	150	2.88	Image
PageBlocks	5393	10	510	9.46	Document
pendigits	6870	16	156	2.27	Image
Pima	768	8	268	34.90	Healthcare
satellite	6435	36	2036	31.64	Astronautics
satimage-2	5803	36	71	1.22	Astronautics
shuttle	49097	9	3511	7.15	Astronautics
skin	245057	3	50859	20.75	Image
smtp	95156	3	30	0.03	Web
SpamBase	4207	57	1679	39.91	Document
speech	3686	400	61	1.65	Linguistics
Stamps	340	9	31	9.12	Document
thyroid	3772	6	93	2.47	Healthcare
vertebral	240	6	30	12.50	Biology
vowels	1456	12	50	3.43	Linguistics
Waveform	3443	21	100	2.90	Physics
WBC	223	9	10	4.48	Healthcare
WDBC	367	30	10	2.72	Healthcare
Wilt	4819	5	257	5.33	Botany
wine	129	13	10	7.75	Chemistry
WPBC	198	33	47	23.74	Healthcare
yeast	1484	8	507	34.16	Biology
CIFAR10	5263	512	263	5.00	Image
FashionMNIST	6315	512	315	5.00	Image
MNIST-C	10000	512	500	5.00	Image
MVTec-AD	5354	512	1258	23.50	Image
SVHN	5208	512	260	5.00	Image
Agnews	10000	768	500	5.00	NLP
Amazon	10000	768	500	5.00	NLP
Imdb	10000	768	500	5.00	NLP
Yelp	10000	768	500	5.00	NLP
20newsgroups	11905	768	591	4.96	NLP

K Differences to Prior Work on PFNs for Tabular Data

There exist applications of PFNs (originally developed by Müller et al. (2022)) that pre-date our proposed FoMo-OD, namely, TabPFN (Hollmann et al., 2023) for supervised classification, LC-PFN (Adriaensen et al., 2024) for learning curve extrapolation, PFN4BO (Müller et al., 2023) for Bayesian optimization, and ForecastPFN (Dooley et al., 2023) for time series forecasting.

Here we highlight the differences of our proposed FoMo-OD from these existing PFNs.

1. **First PFN4OD:** We employ prior-data fitted networks (PFNs) for outlier detection (OD) for the first time.
2. **First large-scale pretrained OD model:** FoMo-OD is the first model for zero-shot OD that is pretrained at large scale on a large collection of (synthetic) datasets, due to the minuscule nature of existing real-world OD benchmark datasets.
3. **New data prior:** Thanks to PFN’s reliance on synthetically generated datasets, we establish a new data prior for OD, specifically for outlier synthesis.
4. **Data transformation for scale:** While drawing samples from a data prior may be relatively fast, pretraining a large foundation model requires many such draws for every step of each epoch. To speed up data synthesis on-the-fly, we are the first to leverage a linear transformation.
5. **Router-based attention for scale:** PFNs ingest the entire training dataset as context for in-context learning at inference time. To accommodate larger datasets at both training (for better generalization) and inference (for large-scale real-world datasets), we leveraged a “bottleneck” architecture for scalable self-attention, and in turn, larger context size.

L Discussion

Summary: We introduced FoMo-OD, **the first foundation model for outlier detection (OD)** on tabular data. FoMo-OD is a prior-data fitted network (PFN), pretrained on a large number of *synthetic* datasets generated from a new data prior for OD, which can infer the posterior predictive distribution for test points in a new dataset in a **zero-shot** fashion where the training data is input as context, capitalizing on *in-context learning*.

Zero-shot OD implies **no additional OD model training or model selection**, given a new OD task. That is a revolution for OD (!), for which algorithm and hyperparameter selection are notoriously-hard *without any labeled data*, and also computationally taxing especially for today’s modern deep OD models with numerous parameters *and* a long list of hyperparameters. What is more, FoMo-OD provides **extremely fast inference** thanks to a mere *single forward pass*, making it amenable for OD on data streams.

Building on the PFN paradigm (Müller et al., 2022), FoMo-OD breaks new ground not only conceptually by abolishing the burden of model training and selection, but also empirically: Against **26** different (both classical and modern) baselines on **57** public benchmark datasets from diverse domains, FoMo-OD performs on par with the top *2nd* baseline, while significantly outperforming the majority of the baselines. Without the need to train any, let alone multiple models for HP tuning, FoMo-OD takes a mere **7.7 ms** per test sample for inference only.

Limitations and Future Directions: FoMo-OD employs a simple straightforward data prior based on GMMs. While it is remarkable to see how far one can go with synthetic data from such a simple prior, future work can design more comprehensive data priors, inclusive of discrete features as well as other possible outlier types. We have also pretrained FoMo-OD solely on synthetic datasets, while future work can augment both synthetic and real-world datasets for pretraining.

Besides the lack of massive real-world datasets for tabular OD, a motivation for a data prior to pretrain purely on synthetic datasets comes from neural scaling laws (Kaplan et al., 2020; Zhai et al., 2022). Interestingly, the scaling laws for large Transformer models have shown that their generalization error tends to drop as a power law with the amount of training data (also, with number of parameters and amount of compute), but

the power law exponent is very small—suggesting that acquiring more colossal real-world datasets would be a slow, if not expensive approach to advancing ML/AI. Others have proposed ways to subset-select smaller, non-redundant “foundation datasets” (Sorscher et al., 2022; Paul et al., 2021), and emphasized the importance of task/dataset diversity in pretraining (Raventós et al., 2024). Arguably, synthetic data from a complex and diverse data prior is a potential gateway to obtaining non-redundant and diverse datasets for pretraining large foundation models like FoMo-OD. On the other hand, designing such a data prior requires a level of domain/prior knowledge.

Another improvement could be scaling up to even larger context (i.e. dataset) size and dimensionality. While FoMo-OD generalizes beyond pretrained context sizes and dimensionality, it is limited to and performs particularly well on downstream datasets of similar nature as our experiments showed. A promising direction for size generalization is using PFNs as extremely fast ensemble components at inference; since “*PFNs are quick enough to be used as ensemble members. The size constraints could therefore be overcome by boosting and bagging techniques*” (Nagler, 2023).

Further, our work focused on unsupervised OD with clean/inlier-only training data. Future work can study the unsupervised OD setting and pretraining with mixed/“contaminated” data in the transductive setting, where the unlabeled test data is the same as training data. In addition, we performed offline evaluation of FoMo-OD on static datasets, while its fast inference lends itself to streaming OD, which future work can explore. Technically, both extensions (unsupervised OD and streaming OD) are straightforward from the implementation perspective.

Our current work is limited to OD for tabular (or point-cloud) data. Our ideas can be extended to other data modalities, such as image, graph, and text outliers, to comprise other domains with critical OD applications such as video surveillance, fraud detection and LLM hallucination detection. To that end, the design of novel inlier/outlier priors would be an open direction. A promising approach here could be the use of pretrained generative models to draw synthesized image/text/etc. datasets for pretraining the PFN, in place of manually-designed data priors.

Finally, our quest here has been mainly experimental. Theoretically understanding why these models work as well as they do and investigating their failure cases are important yet open questions. As empirical future work, one could systematically stress-test FoMo-OD using synthetically generated test datasets that contain known outlier types and inlier distributions distinct from those used during pretraining, while varying the degree of out-of-distribution characteristics in a controlled manner.

As the first foundation model for OD, FoMo-OD inspires many promising directions for future research that could lead to fruition for additional practical applications.

M Reproducibility Statement

We expect that the disruptive nature of FoMo-OD will trigger future innovations in the OD literature, as well as a widespread adoption by practitioners thanks to its key desirable properties. To foster future research and accessibility in practice, we make all resources (our codebase used for prior data synthesis, data transformation, and pretraining as well as our pretrained model checkpoints) publicly available at <https://github.com/A-Chicharito-S/FoMo-OD>. Full implementation details are provided in Appendix C.