

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

Anonymous Authors¹

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. Having been pre-trained on synthetic data, FoMo-0D can directly predict the (outlier/inlier) label of test samples without parameter fine-tuning—*requiring 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://anonymous.4open.science/r/PFN40D>.

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—mainly because compared to classical methods, deep OD models (Pang et al., 2021) have many more HPs that detection performance is sensitive to (Ding et al., 2022), rendering 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 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.

In this paper, we capitalize on these ideas and introduce FoMo-0D: a prior-data fitted Foundation Model 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.

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

¹Anonymous Institution, Anonymous City, Anonymous Region, Anonymous Country. Correspondence to: Anonymous Author <anon.email@domain.com>.

Preliminary work. Under review by the International Conference on Machine Learning (ICML). Do not distribute.

Table 1: p -values of the one-sided Wilcoxon signed rank test, comparing FoMo-OD (with $D = 100$) to **top 10 baselines** with default hyperparameters (HPs), and **top 4^{avg}** baselines³ 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. FoMo-OD shows **no statistically significant difference from the 2nd best model** (k NN, w/ $p = 0.106$) over All datasets, 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). Rank, avg.'ed over all 57 datasets by AUROC. (setting: $D = 100$, $R = 500$, train/inference context size=5K, w/ quantile transform) (See Tables 16.1&16.2 for full results.)

	FoMo-OD	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
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

In designing FoMo-OD, 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), 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-OD achieves remarkable performance: As shown in Table 1, FoMo-OD, 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 16.1&16.2 for full results.) The results remain consistent on (46) benchmarks with $d \leq 500$ dimensions. FoMo-OD is also competitive across all (57) datasets, effectively generalizing beyond its pre-training distributions, with no statistically significant difference from the 2nd best baseline. Further, FoMo-OD takes a mere 7.7 ms to infer a test sample on average with no extra training or tuning overhead on the new dataset.

2. Problem and Preliminaries

Semi-supervised Outlier Detection: We focus on semi-supervised OD. Formally, let $\mathcal{D}_{\text{in}} = \{(\mathbf{x}_1, y_1) \dots, (\mathbf{x}_n, y_n)\}$ denote the input data containing only inliers $\mathbf{x}_i \in \mathbb{R}^d$, where $y_i = 0 \forall i \in [n]$, and $\mathcal{D}_{\text{test}}$ depicts the test data comprising both inliers and outliers. The task is to assign labels to $\mathbf{x}_i \in \mathcal{D}_{\text{test}}$ given the inlier-only input $\mathcal{D}_{\text{in}}$.

PFNs: (Müller et al., 2023) approximate the posterior predictive distribution (PPD) of a test point using training data as context, formulated as $p(\cdot | \mathbf{x}_{\text{test}}, \mathcal{D}_{\text{train}})$. A Transformer is pre-trained on synthetic datasets from predefined priors and used at inference for zero-shot prediction with test and

context samples. More details in Appendix C.

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 as follows.

$$\mathcal{L} = \mathbb{E}_{(\{\mathbf{x}_{\text{test}}, y_{\text{test}}\}) \cup \mathcal{D}_{\text{train}} \sim p(\mathcal{D})} [-\log q_{\theta}(y_{\text{test}} | \mathbf{x}_{\text{test}}, \mathcal{D}_{\text{train}})] \quad (1)$$

Inference on real-world data. At inference, a fresh real-world dataset $\mathcal{D}_{\text{train}}$ and test instance $\mathbf{x}_{\text{test}}$ are fed into the 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).

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

3.1. Designing a Data Prior for Outlier Detection

We design a new data prior from which we simulate numerous OD datasets for pre-training FoMo-OD.

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 diagonal² 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 gener-

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.

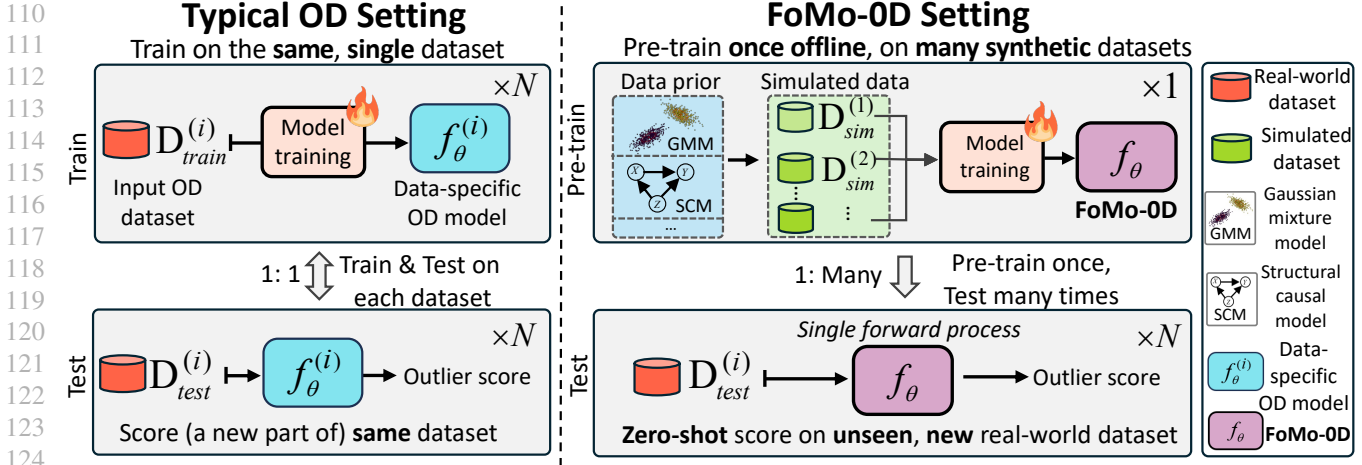


Figure 1: (best in color) Comparison of typical OD vs. the FoMo-OD settings. Given a new unlabeled OD dataset, FoMo-OD not only eliminates the need for model (parameter) training, but most importantly, also abolishes the onerous task of unsupervised model selection (algorithm and hyperparameters).

ate 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 E.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 B.

3.2. (Pre)Training and Inference

Model (Pre)Training (Once, Offline): In the synthetic prior-data fitting phase, we first randomly draw a hypothesis (i.e. GMM configuration) uniformly at random, i.e., $\phi = \{d \in [D], m \in [M], \{\mu_j\}_{j=1}^m \in [-5, 5]^d, \{\Sigma_j\}_{j=1}^m; \text{diag}(\Sigma_j) \in [-5, 5]^d\}$, 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-OD’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. (1)). 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 $n_U = 5,000$.

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

3.3. Architecture and Scalability

FoMo-OD is based on Transformer (Vaswani et al., 2017) model, with (1) self-attention among all the training samples and only cross-attention from test samples to the training samples; (2) “router attention mechanism” Zhang & Yan (2023) to reduce the complexity of self-attention to linear (As shown in Figure 2), thus extending the context length. We further propose to scale up (pre)training data synthesis with linear transforms, more details in Appendix D.

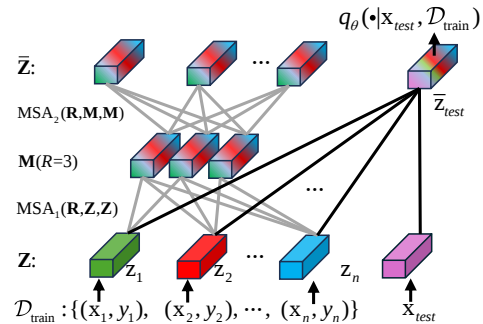


Figure 2: FoMo-OD architecture employs the “router mechanism” for scalable attention.

4. Experiments

4.1. Setup

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 $\{\mu_j\}_{j=1}^m$ (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 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 19 in Appendix N). Following Livernoche et al. (2024), we use 5 train/test splits of each dataset via different seeds and report mean performance and standard deviation.

Baselines: We compare FoMo-0D against 26 baselines, from classical/shallow methods to modern/deep models, imported from one of the latest papers that proposed the SOTA diffusion-based OD model, DTE (Livernoche et al., 2024). We refer to the original paper for more details.

Model Implementation: We train FoMo-0D for 200,000 steps with a batch size of 8 datasets (i.e., 1,600,000 synthetically generated datasets), which 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 (e.g., the training algorithm, model architecture, data synthesis) are in Appendix F.

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 significant at $\alpha = 0.05$ following convention.

Hyperparameters (HPs): 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**³ 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 HPs selected at random, in the absence of any other prior knowledge. We annotate the method name with ^{avg} for the version with performance averaged over varying HPs. The list of HP values for each top baseline is detailed in Appendix G.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 M**. 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 $d \leq 500$. We present further results, ablation analyses, generalization analyses in Appendix I, J, K.

5. Conclusion

This work introduced FoMo-0D, **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-0D breaks new ground by fully abolishing the notoriously-hard model selection task for unsupervised OD (see Impact Statement). Further, FoMo-0D 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-0D 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 P.

³ To rank the 26 baselines, we compute the 26×26 p -values of the pairwise Wilcoxon signed rank test (see Appendix Figure 23), and order them by their mean p -value against other baselines.

References

- Achiam, J., Adler, S., Agarwal, S., Ahmad, L., Akkaya, I., Aleman, F. L., Almeida, D., Altenschmidt, J., Altman, S., Anadkat, S., et al. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*, 2023.
- Adriaensen, S., Rakotoarison, H., Müller, S., and Hutter, F. Efficient bayesian learning curve extrapolation using prior-data fitted networks. *Advances in Neural Information Processing Systems*, 36, 2024.
- Aggarwal, C. C. *Outlier Analysis*. Springer, 2013.
- Aggarwal, C. C. and Sathe, S. Theoretical foundations and algorithms for outlier ensembles. *Acm sigkdd explorations newsletter*, 17(1):24–47, 2015.
- Aggarwal, C. C. and Sathe, S. *Outlier Ensembles - An Introduction*. Springer, 2017.
- Akçay, S., Atapour-Abarghouei, A., and Breckon, T. P. Ganomaly: Semi-supervised anomaly detection via adversarial training. In *ACCV*, pp. 622–637. Springer, 2019.
- Alayrac, J.-B., Donahue, J., Luc, P., Miech, A., Barr, I., Hasson, Y., Lenc, K., Mensch, A., Millican, K., Reynolds, M., et al. Flamingo: a visual language model for few-shot learning. *Advances in neural information processing systems*, 35:23716–23736, 2022.
- Bergman, L. and Hoshen, Y. Classification-based anomaly detection for general data. In *International Conference on Learning Representations*, 2020.
- Breunig, M. M., Kriegel, H.-P., Ng, R. T., and Sander, J. Lof: Identifying density-based local outliers. In *International Conference on Management of Data*, 2000.
- Brown, T., Mann, B., Ryder, N., Subbiah, M., Kaplan, J. D., Dhariwal, P., Neelakantan, A., Shyam, P., Sastry, G., Askell, A., Agarwal, S., Herbert-Voss, A., Krueger, G., Henighan, T., Child, R., Ramesh, A., Ziegler, D., Wu, J., Winter, C., Hesse, C., Chen, M., Sigler, E., Litwin, M., Gray, S., Chess, B., Clark, J., Berner, C., McCandlish, S., Radford, A., Sutskever, I., and Amodei, D. Language models are few-shot learners. In *Advances in Neural Information Processing Systems*, volume 33, pp. 1877–1901, 2020.
- Campos, G. O., Zimek, A., Sander, J., Campello, R. J., Mícenková, B., Schubert, E., Assent, I., and Houle, M. E. On the evaluation of unsupervised outlier detection: measures, datasets, and an empirical study. *Data mining and knowledge discovery*, 30:891–927, 2016.
- Casella, G. and Berger, R. *Statistical inference*. CRC Press, 2024.
- Chalapathy, R. and Chawla, S. Deep learning for anomaly detection: A survey. *arXiv preprint arXiv:1901.03407*, 2019.
- Demšar, J. Statistical comparisons of classifiers over multiple data sets. *The Journal of Machine learning research*, 7:1–30, 2006.
- Ding, X., Zhao, L., and Akoglu, L. Hyperparameter sensitivity in deep outlier detection: Analysis and a scalable hyper-ensemble solution. *Advances in Neural Information Processing Systems*, 35:9603–9616, 2022.
- Ding, X., Zhao, Y., and Akoglu, L. Fast unsupervised deep outlier model selection with hypernetworks. *ACM SIGKDD*, 2024.
- Dolan, E. D. and Moré, J. J. Benchmarking optimization software with performance profiles. *Mathematical programming*, 91:201–213, 2002.
- Dooley, S., Khurana, G. S., Mohapatra, C., Naidu, S. V., and White, C. ForecastPFN: Synthetically-trained zero-shot forecasting. In *Thirty-seventh Conference on Neural Information Processing Systems*, 2023.
- Du, X., Sun, Y., Zhu, J., and Li, Y. Dream the impossible: Outlier imagination with diffusion models. *Advances in Neural Information Processing Systems*, 36, 2024.
- Esmaeilpour, S., Liu, B., Robertson, E., and Shu, L. 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.
- Feuer, B., Cohen, N., and Hegde, C. Scaling tabPFN: Sketching and feature selection for tabular prior-data fitted networks. In *NeurIPS 2023 Second Table Representation Learning Workshop*, 2023.
- Feuer, B., Schirrmeister, R. T., Cherepanova, V., Hegde, C., Hutter, F., Goldblum, M., Cohen, N., and White, C. Tunetables: Context optimization for scalable prior-data fitted networks, 2024.
- Garg, S., Tsipras, D., Liang, P. S., and Valiant, G. What can transformers learn in-context? a case study of simple function classes. *Advances in Neural Information Processing Systems*, 2022.
- Goldstein, M. and Dengel, A. Histogram-based outlier score (hbos): A fast unsupervised anomaly detection algorithm. *KI-2012: poster and demo track*, 1:59–63, 2012.
- Goldstein, M. and Uchida, S. A comparative evaluation of unsupervised anomaly detection algorithms for multivariate data. *PloS one*, 11(4), 2016.

- Goyal, S., Raghunathan, A., Jain, M., Simhadri, H., and Jain, P. Drocc: Deep robust one-class classification. In *Proceedings of the 37th International Conference on Machine Learning, ICML'20*. JMLR.org, 2020.
- Gu, Z., Zhu, B., Zhu, G., Chen, Y., Tang, M., and Wang, J. AnomalyGPT: Detecting industrial anomalies using large vision-language models. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 38, pp. 1932–1940, 2024.
- Gut, A. *An Intermediate Course in Probability*. Springer Texts in Statistics. Springer New York, 2009. ISBN 9781441901620.
- Han, S., Hu, X., Huang, H., Jiang, M., and Zhao, Y. Ad-bench: Anomaly detection benchmark. *Advances in Neural Information Processing Systems*, 35, 2022.
- He, H., Zhang, J., Chen, H., Chen, X., Li, Z., Chen, X., Wang, Y., Wang, C., and Xie, L. A diffusion-based framework for multi-class anomaly detection. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 38, pp. 8472–8480, 2024.
- He, K., Zhang, X., Ren, S., and Sun, J. Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pp. 770–778, 2016.
- Ho, J., Jain, A., and Abbeel, P. Denoising diffusion probabilistic models. *Advances in Neural Information Processing Systems*, 33:6840–6851, 2020.
- Hojjati, H., Ho, T. K. K., and Armanfard, N. Self-supervised anomaly detection: A survey and outlook. *arXiv preprint arXiv:2205.05173*, 2022.
- Hollmann, N., Müller, S., Eggensperger, K., and Hutter, F. TabPFN: A transformer that solves small tabular classification problems in a second. In *The Eleventh International Conference on Learning Representations*, 2023.
- Huber-Carol, C., Balakrishnan, N., Nikulin, M., and Mesbah, M. *Goodness-of-fit tests and model validity*. Springer Science & Business Media, 2012.
- Jeong, J., Zou, Y., Kim, T., Zhang, D., Ravichandran, A., and Dabeer, O. 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.
- Kaplan, J., McCandlish, S., Henighan, T., Brown, T. B., Chess, B., Child, R., Gray, S., Radford, A., Wu, J., and Amodei, D. Scaling laws for neural language models. *arXiv preprint arXiv:2001.08361*, 2020.
- Kingma, D. P. Auto-encoding variational bayes. *arXiv preprint arXiv:1312.6114*, 2013.
- Kingma, D. P. and Ba, J. Adam: A method for stochastic optimization, 2017.
- Li, A., Qiu, C., Kloft, M., Smyth, P., Rudolph, M., and Mandt, S. Zero-shot anomaly detection via batch normalization. In *Thirty-seventh Conference on Neural Information Processing Systems*, 2023.
- Li, A., Zhao, Y., Qiu, C., Kloft, M., Smyth, P., Rudolph, M., and Mandt, S. Anomaly detection of tabular data using LLMs. *arXiv preprint arXiv:2406.16308*, 2024.
- Liu, F. T., Ting, K. M., and Zhou, Z.-H. Isolation forest. In *2008 Eighth IEEE International Conference on Data Mining*, pp. 413–422, 2008.
- Liu, H., Li, C., Wu, Q., and Lee, Y. J. Visual instruction tuning. *Advances in neural information processing systems*, 36, 2024.
- Livernoche, V., Jain, V., Hezaveh, Y., and Ravanbakhsh, S. On diffusion modeling for anomaly detection. In *ICLR*, 2024.
- Liznerski, P., Ruff, L., Vandermeulen, R. A., Franks, B. J., Müller, K.-R., and Kloft, M. Exposing outlier exposure: What can be learned from few, one, and zero outlier images. *arXiv preprint arXiv:2205.11474*, 2022.
- Ma, J., Thomas, V., Yu, G., and Caterini, A. L. In-context data distillation with tabpfn. *CoRR*, abs/2402.06971, 2024.
- Ma, M. Q., Zhao, Y., Zhang, X., and Akoglu, L. The need for unsupervised outlier model selection: A review and evaluation of internal evaluation strategies. *ACM SIGKDD Explorations Newsletter*, 25(1):19–35, 2023.
- Müller, S., Feurer, M., Hollmann, N., and Hutter, F. PFNs4BO: in-context learning for bayesian optimization. In *International Conference on Machine Learning*, 2023.
- Müller, S., Hollmann, N., Pineda-Arango, S., Grabocka, J., and Hutter, F. Transformers can do bayesian inference. *ICLR*, 2022.
- Nagler, T. Statistical foundations of prior-data fitted networks. In *ICML*, volume 202 of *Proceedings of Machine Learning Research*. PMLR, 2023.
- Pang, G., Shen, C., Cao, L., and Hengel, A. V. D. Deep learning for anomaly detection: A review. *ACM computing surveys (CSUR)*, 54(2):1–38, 2021.

- Paul, M., Ganguli, S., and Dziugaite, G. K. Deep learning on a data diet: Finding important examples early in training. *Advances in neural information processing systems*, 34: 20596–20607, 2021.
- Radford, A., Kim, J. W., Hallacy, C., Ramesh, A., Goh, G., Agarwal, S., Sastry, G., Askell, A., Mishkin, P., Clark, J., et al. Learning transferable visual models from natural language supervision. In *International conference on machine learning*, 2021.
- Ramaswamy, S., Rastogi, R., and Shim, K. Efficient algorithms for mining outliers from large data sets. In *ACM SIGMOD international conference on management of data*, 2000.
- Ramesh, A., Pavlov, M., Goh, G., Gray, S., Voss, C., Radford, A., Chen, M., and Sutskever, I. Zero-shot text-to-image generation. In *International conference on machine learning*, pp. 8821–8831. Pmlr, 2021.
- Raventós, A., Paul, M., Chen, F., and Ganguli, S. Pretraining task diversity and the emergence of non-bayesian in-context learning for regression. *Advances in Neural Information Processing Systems*, 2024.
- Rezende, D. and Mohamed, S. 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.
- Ruff, L., Vandermeulen, R., Goernitz, N., Deecke, L., Siddiqui, S. A., Binder, A., Müller, E., and Kloft, M. Deep one-class classification. In *International Conference on Machine Learning*, pp. 4393–4402, 2018.
- Ruff, L., Kauffmann, J. R., Vandermeulen, R. A., Montavon, G., Samek, W., Kloft, M., Dietterich, T. G., and Müller, K.-R. A unifying review of deep and shallow anomaly detection. *Proceedings of the IEEE*, 109(5):756–795, 2021.
- Shenkar, T. and Wolf, L. Anomaly detection for tabular data with internal contrastive learning. In *International Conference on Learning Representations*, 2022.
- Sorscher, B., Geirhos, R., Shekhar, S., Ganguli, S., and Morcos, A. Beyond neural scaling laws: beating power law scaling via data pruning. *Advances in Neural Information Processing Systems*, 35:19523–19536, 2022.
- Steinbuss, G. and Böhm, K. Benchmarking unsupervised outlier detection with realistic synthetic data. *ACM Transactions on Knowledge Discovery from Data (TKDD)*, 15(4):1–20, 2021.
- Touvron, H., Lavril, T., Izacard, G., Martinet, X., Lachaux, M.-A., Lacroix, T., Rozière, B., Goyal, N., Hambro, E., Azhar, F., Rodriguez, A., Joulin, A., Grave, E., and Lample, G. Llama: Open and efficient foundation language models, 2023. URL <https://arxiv.org/abs/2302.13971>.
- Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, L., and Polosukhin, I. Attention is all you need. In *NIPS*, pp. 5998–6008, 2017.
- Xie, S. M., Raghunathan, A., Liang, P., and Ma, T. An explanation of in-context learning as implicit bayesian inference. *arXiv preprint arXiv:2111.02080*, 2021.
- Xu, H., Wang, Y., Wei, J., Jian, S., Li, Y., and Liu, N. 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.
- Yoo, J., Zhao, T., and Akoglu, L. 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.
- Yoon, S., Jin, Y.-U., Noh, Y.-K., and Park, F. Energy-based models for anomaly detection: A manifold diffusion recovery approach. *Advances in Neural Information Processing Systems*, 36:49445–49466, 2023.
- Zhai, X., Kolesnikov, A., Houlsby, N., and Beyer, L. Scaling vision transformers. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 2022.
- Zhang, J., Yang, J., Wang, P., Wang, H., Lin, Y., Zhang, H., Sun, Y., Du, X., Zhou, K., Zhang, W., et al. Openood v1.5: Enhanced benchmark for out-of-distribution detection. *arXiv preprint arXiv:2306.09301*, 2023.
- Zhang, Y. and Yan, J. Crossformer: Transformer utilizing cross-dimension dependency for multivariate time series forecasting. In *ICLR*, 2023.
- Zhao, Y. and Akoglu, L. Toward unsupervised outlier model selection. In *International Conference on Automated Machine Learning (AutoML)*, 2024.
- Zhao, Y., Rossi, R., and Akoglu, L. Automatic unsupervised outlier model selection. *Advances in Neural Information Processing Systems*, 34:4489–4502, 2021.
- Zhao, Y., Zhang, S., and Akoglu, L. Toward unsupervised outlier model selection. In *2022 IEEE International Conference on Data Mining (ICDM)*, pp. 773–782. IEEE, 2022.

Zhou, Q., Pang, G., Tian, Y., He, S., and Chen, J. 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>.

Zong, B., Song, Q., Min, M. R., Cheng, W., Lumezanu, C., Ki Cho, D., and Chen, H. Deep autoencoding gaussian mixture model for unsupervised anomaly detection. In *International Conference on Learning Representations*, 2018.

Table of Contents

We detail the contents in the appendix below.

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

A. 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 (an extremely fast ensemble of randomized trees).

More recent UOMS approaches go beyond heuristic measures and instead design scalable hyperensembles (Ding et al., 2022; 2024), as well as take advantage of meta-learning on historical real-world OD datasets (Zhao et al., 2021; 2022; Zhao & Akoglu, 2024). These approaches demonstrate the value of learning from many other OD datasets, and transfer these learnings to a new dataset. While sharing the same spirit on learning from a large collection of (in our case, simulated) datasets, our FoMo-OD differs from these prior art in a key aspect; FoMo-OD is *not* a model selection technique, but rather, a foundation model that abolishes model training and selection altogether—unlocking 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) 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. As training data is passed as context to PFN, others proposed scaling solutions to enable training on larger pretraining datasets for better generalization (Ma et al., 2024; Feuer et al., 2023; 2024).

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 O 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 (Esmailpour 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-0D employs and pretrains PFNs at a much larger scale with rigorous evaluation on a much larger testbed.

B. Illustration of synthetic data in 2-d

We visualize our synthetic data in Figure 3, 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.

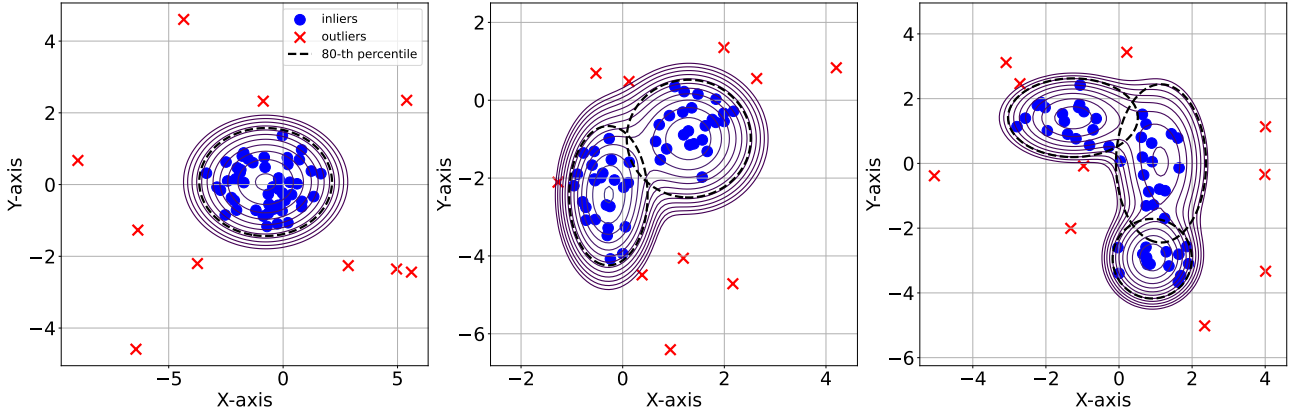


Figure 3: Illustration of synthetic data in 2D with 80th percentile as the criterion.

C. 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. 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 :

$$p(y_{\text{test}} | \mathbf{x}_{\text{test}}, \mathcal{D}_{\text{train}}) = \int_{\Phi} p(y_{\text{test}} | \mathbf{x}_{\text{test}}, \phi) p(\mathcal{D}_{\text{train}} | \phi) p(\phi) d\phi, \quad (2)$$

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 as follows.

$$\mathcal{L} = \mathbb{E}_{\{(\mathbf{x}_{\text{test}}, y_{\text{test}})\} \cup \mathcal{D}_{\text{train}} \sim p(\mathcal{D})} [-\log q_{\theta}(y_{\text{test}} | \mathbf{x}_{\text{test}}, \mathcal{D}_{\text{train}})] . \quad (3)$$

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).

D. 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 token, 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

$$\mathbf{z}_i = \sum_{j=1}^n \text{softmax}(\{\langle \mathbf{q}_i, \mathbf{k}_{j'} \rangle\}_{j'=1}^n)_j \cdot \mathbf{v}_j. \quad (4)$$

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 & Sathé, 2015; Campos et al., 2016; Goldstein & Uchida, 2016; Ding et al., 2022). We present analyses on sample-to-sample attention in Appendix H.

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” or representatives, which gather information from all n samples and then distribute the information back to the n output embeddings—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.

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

$$\mathcal{M} = \text{MSA}_1(\mathbf{R}, \mathbf{Z}, \mathbf{Z}), \quad (5)$$

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:

$$\hat{\mathbf{Z}} = \text{MSA}_2(\mathbf{Z}, \mathcal{M}, \mathcal{M}) . \quad (6)$$

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, 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 6). 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 E.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 E.

Lemma D.1. *Linear transform T with invertible $\mathbf{W}$ on $\mathcal{G}_m^d$ preserves Mahalanobis distances.*

Lemma D.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).

E. Linear Transform for Scalable GMM Data Synthesis

E.1. Definitions

Definition E.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:

$$p(\mathbf{x}) = \sum_{j=1}^m w_j \cdot g(\mathbf{x} | \boldsymbol{\mu}_j, \boldsymbol{\Sigma}_j) , \quad (7)$$

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

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 $\boldsymbol{\Sigma}_j$'s are favored in our data prior.) In addition, Mahalanobis distance for labeling inliers/outliers also requires the inverse.

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.

$\mathbf{x}^T \Sigma_i \mathbf{x} \geq 0$, for all $\mathbf{x} \in \mathbb{R}^d$.

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

$$T(\mathbf{x}) = \mathbf{W}\mathbf{x} + \mathbf{b}, \quad (8)$$

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

$$\text{dist}_M(\mathbf{x}) = \sqrt{(\mathbf{x} - \boldsymbol{\mu})^T \Sigma^{-1} (\mathbf{x} - \boldsymbol{\mu})}. \quad (9)$$

Definition E.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.

E.2. Properties

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

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

$$\text{dist}_M^2(\mathbf{x}) = \mathbf{z}^T \mathbf{z} = \sum_{i=1}^d z_i^2 \quad (10)$$

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

E.3. Lemmas

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

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

$$\text{dist}_M(T(\mathbf{x})) = \sqrt{(T(\mathbf{x}) - (\mathbf{W}\boldsymbol{\mu}_j + \mathbf{b}))^T (\mathbf{W}\Sigma_j \mathbf{W}^T)^{-1} (T(\mathbf{x}) - (\mathbf{W}\boldsymbol{\mu}_j + \mathbf{b}))} \quad (11)$$

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

$$= \sqrt{(\mathbf{x} - \boldsymbol{\mu}_j)^T \mathbf{W}^T (\mathbf{W}^T)^{-1} \Sigma_j^{-1} \mathbf{W} (\mathbf{x} - \boldsymbol{\mu}_j)} \quad (13)$$

$$= \sqrt{(\mathbf{x} - \boldsymbol{\mu}_j)^T \Sigma_j^{-1} (\mathbf{x} - \boldsymbol{\mu}_j)} = \text{dist}_M(\mathbf{x}). \quad (14)$$

□

Lemma E.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 $X \sim \chi_d^2$:

$$\text{Prob}(X \leq \chi_d^2(\alpha)) = \frac{\alpha}{100}. \quad (15)$$

Based on Property E.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 $\mathcal{N}(\boldsymbol{\mu}_j, \Sigma_j)$, 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 E.7), then we conclude that the linear transform T with invertible $\mathbf{W}$ preserves the percentiles of the GMM. □

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_θ .
- 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 $D_j := \mathcal{D}_{\text{train}} \cup \{(\mathbf{x}_k, y_k)\}_{k=1}^{|\mathcal{D}_{\text{test}}|} \sim p(\mathcal{D})$.
- 10: $\mathcal{C}_q \leftarrow \mathcal{C}_q + [D_j]$
- 11: **else**
- 12: $j \leftarrow j \bmod q$
- 13: $D_j \leftarrow T(\mathcal{C}_q[j])$
- 14: **end if**
- 15: **Step 2:** compute stochastic loss approximation $\bar{\ell}_\theta = \sum_{k=1}^{|\mathcal{D}_{\text{test}}|} (-\log q_\theta(y_k | \mathbf{x}_k, \mathcal{D}_{\text{train}}))$.
- 16: **Step 3:** update parameters θ with stochastic gradient descent on $\nabla_\theta \bar{\ell}_\theta$.
- 17: $\mathcal{B}_i \leftarrow \mathcal{B}_i + [D_j]$
- 18: **end for**
- 19: $\mathcal{C}_E \leftarrow \mathcal{C}_E + [\mathcal{B}_i]$
- 20: **else**
- 21: $i \leftarrow i \bmod P$
- 22: $\mathcal{B}_i \leftarrow \mathcal{C}_E[i]$
- 23: **for** $j \leftarrow 1$ **to** Q **do**
- 24: $D_j \leftarrow T(\mathcal{B}_i[j])$
- 25: Perform Step 2 and Step 3
- 26: **end for**
- 27: **end if**
- 28: **end for**

F. Implementation details

F.1. Hardware

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

F.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 For an input with d features, we follow Müller et al. (2022) and deal with $d < D$ by rescaling the input with $\frac{D}{d}$ and padding the features to size D with 0, and randomly sample D features out of d if $d > D$. In addition, FoMo-0D uses context size of 5K at inference, where we randomly sample $(5K-1)$ points as $\mathcal{D}_{\text{train}}$ from datasets with $n > 5K$ for each test sample $\mathbf{x} \in \mathcal{D}_{\text{test}}$.

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

Training loss In Figure 4, 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

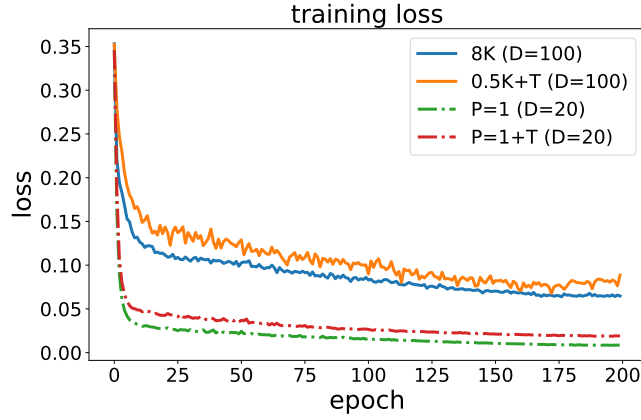


Figure 4: (best in color) Training loss of FoMo-0D ($D = 100$) with 8K unique datasets/epoch (in blue) and using 0.5K unique + 7.5K transformed datasets/epoch (in orange), and FoMo-0D ($D = 20$) with $P = 1$ (in green) and $P = 1$ with transformation (in red) over 200 epochs.

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

Figure 5 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.

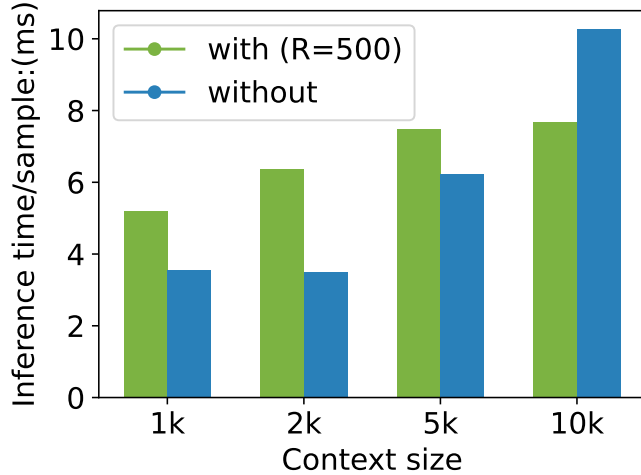


Figure 5: Inference time of FoMo-0D on GPU with vs. w/out router-based attention under varying context size.

G. Detailed Experiment Setup

G.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 D.⁶ 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.

G.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 19). 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 semi-supervised OD setup 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.

G.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 (Akcay 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.).

G.4. Hyperparameters for Baselines

Table 2 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 2: 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	$\text{learning_rate} \in \{10^{-1}, 10^{-2}, 10^{-3}, 10^{-4}, 10^{-5}\}$
DTE-C	$k \in \{5, 10, 20, 40, 50\}$

⁶It 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 $\mathbf{b} \in [-1, 1]^d$.

G.5. Ranking the 26 baselines

Figure 23 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.

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

Figure 24, 25, 26, 27 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, $learning_rate$ (lr) $\in \{10^{-1}, 10^{-2}, 10^{-3}, 10^{-4}, 10^{-5}\}$; and for others, $\#nearest_neighbors$ $k \in \{5, 10, 20, 40, 50\}$). We find that for ICL, $lr = 10^{-3}$ 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 kNN^{avg} models across HP configurations perform subpar

G.7. Sampling time of d -dimensional GMM

Figure 6 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.

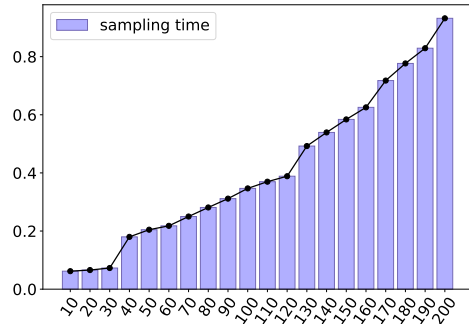


Figure 6: Sampling time (in seconds) of 10,000 points from GMMs with varying number of dimensions.

H. 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 7, 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 8. For a specific outlier, there is a similar trend of attending to the center of a Gaussian (as shown in Figure 7), 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),

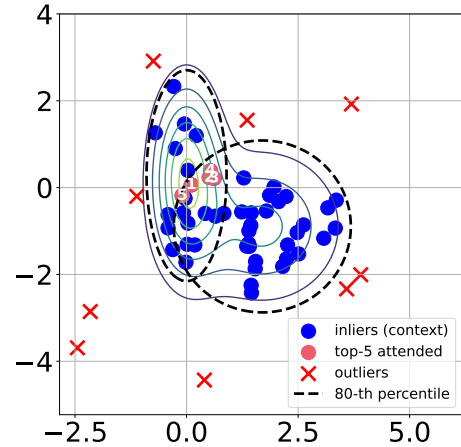


Figure 7: Top-5 attended inliers (all 50 inliers and only part of the outliers are shown for better visualization).

suggesting FoMo-OD is incorporating both boundary and nearest neighbor information dynamically for each outlier.

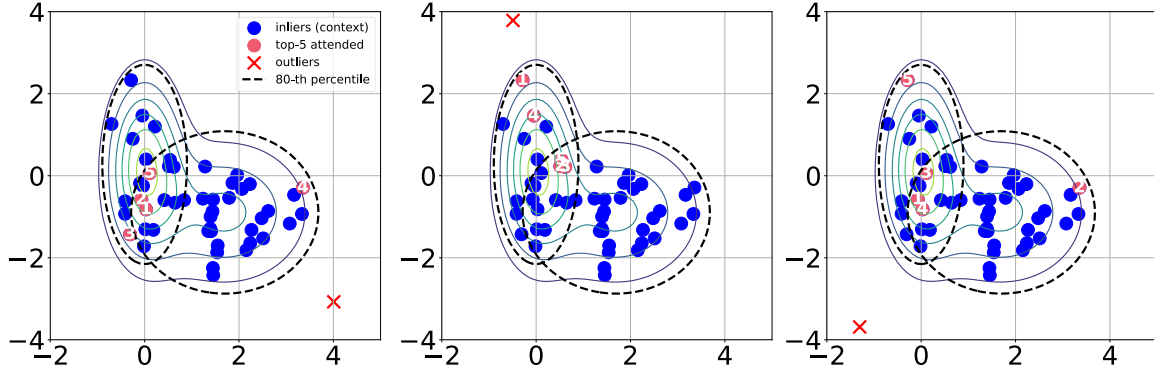


Figure 8: Top-5 attended inliers of 3 outliers at different positions of the GMM

I. Further Results

Detection performance: Table 3 shows similar results for FoMo-OD 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-OD pre-trained purely on synthetic datasets from a simple prior in small dimensions, these results showcase the prowess of PFNs for OD.

Table 3: p -values of the one-sided Wilcoxon signed rank test, comparing FoMo-OD (w/ $D = 20$) to **top 10** baselines with default HPs, and **top 4^{avg}** baselines³ 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-OD 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). 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)

	FoMo-OD	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

Figure 9 shows the distribution of ranks across datasets for all models. While p -values are the most statistically conclusive, FoMo-OD achieves a relatively small average rank with notably lower ranks across datasets compared to the majority of the baselines. Appendix L presents another comparison between detectors through performance profile plots (Dolan & Moré, 2002).

Running time: Table 4 presents the total training time and the average inference time per test sample as measured on the largest dataset for FoMo-OD and the top-3 baselines. Given a new dataset, FoMo-OD bypasses model training (and HP tuning) and directly performs inference, with an average of 7.7 ms per sample (see Appendix Figure 5). In comparison, all baseline methods need to train on each individual dataset preceding inference. This training time can be high for deep learning based models like ICL, and further compounded with training multiple models for hyperparameter tuning purposes. Even for non-parametric and/or shallow models like k NN and DTE-NP (which queries k nearest neighbors), the training involves various data pre-processing steps such as constructing a tree-like data structure for fast (often approximate) k NN distance querying.

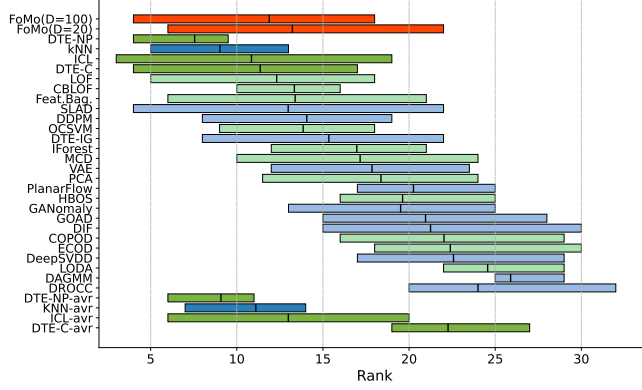


Figure 9: (best in color) Rank (w.r.t. AUROC, lower is better) distribution across all **57** real-world datasets shown via boxplots for (from top to bottom) FoMo-OD in **red**, all **26** baselines ordered by mean p -value³ (shallow and deep baselines in **green** and **blue**), and **top 4** baselines’ avg variants.

Table 4: Train-time and Inference-time (in milliseconds) of FoMo-OD and the top-3³ baselines (w/ *default* HPs, *excluding* the time for model selection/hyperparameter optimization) on our largest dataset `donors` (see Appendix Table 19). FoMo-OD skips any model training or fine-tuning and takes a mere forward pass for inference out-of-the-box.

Method	FoMo-OD	DTE-NP	kNN	ICL
Train-t (total)	none/0-shot	56.83	1433.74	186461.48
Infer-t (per sample)	7.7	0.76	0.17	0.01

J. Ablation Analyses

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

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

J.1. Effect of pretraining dimensionality D

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

We start by comparing FoMo-OD 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 4. 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 3 ($D = 20$) w.r.t. p -values over **All** datasets, we find that FoMo-OD at larger scale does better, where **all** p -values are larger for $D = 100$ than $D = 20$. We find that FoMo-OD 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 12.

J.2. Effect of routers on cost

What is the running time and memory cost of FoMo-OD with & w/out router-based attention?

Figure 10(left) shows the average inference time per test sample, comparing FoMo-OD using a router-based attention mechanism with $R = 500$ routers (in green) versus FoMo-OD 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.⁷

⁷Note that the inference time is reported on CPUs to show scalability. On GPUs, w/ 5K context size, see Appendix Figure 5, where

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

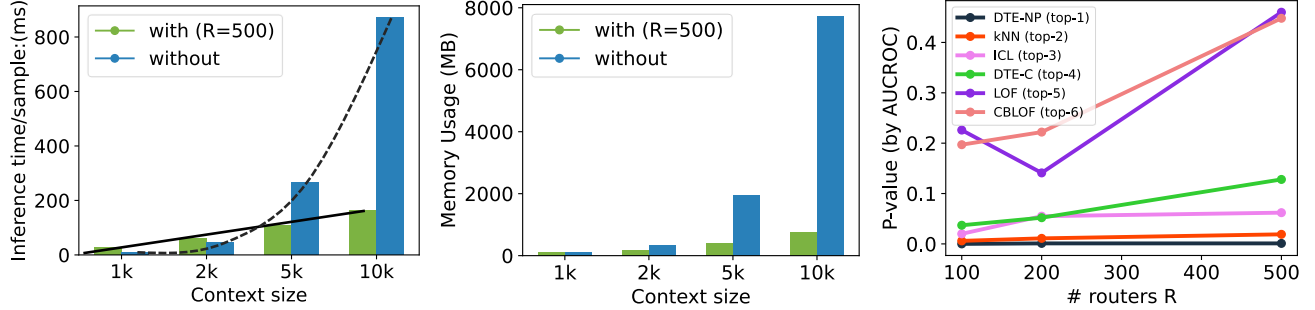


Figure 10: FoMo-0D w/ router mechanism saves time and memory while more #routers perform better, offering a cost-performance trade-off: (left) inference-time (ms) per sample and (middle) memory cost (MB) with & w/out routers by varying context size; (right) performance (based on p -value against top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

J.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.

J.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.

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 (5) and (6).

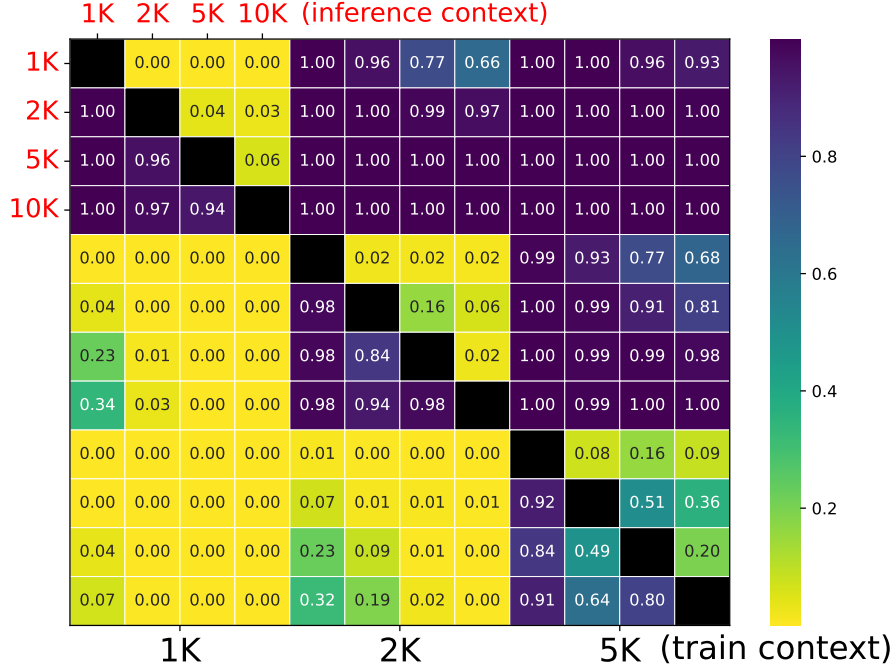


Figure 11: p -values of the pairwise Wilcoxon signed rank test between models (larger p implies col-method is better than row-method) w/ different context sizes for **training** (1K/2K/5K, 1st/2nd/3rd four grids, in **black**) and **inference** (1K/2K/5K/10K, every 1st/2nd/3rd/4th grid, in **red**):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

J.5. Effect of number of unique datasets

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

Next we study the effect of dataset *reuse at epoch level* (w/out transformation) on performance as presented in Section D. 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 $d \leq 50$.

J.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

J.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 4 and Appendix Figure 5).

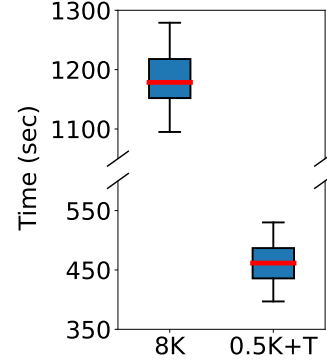


Figure 12: Runtime/epoch distribution over 100 epochs for FoMo-0D ($D=100$) with (left) $P=100$, i.e. 8K unique datasets/epoch vs. (right) 0.5K unique+7.5K transformed datasets/epoch.

J.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.

J.9. Effect of applying quantile transform on benchmark datasets

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

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 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.

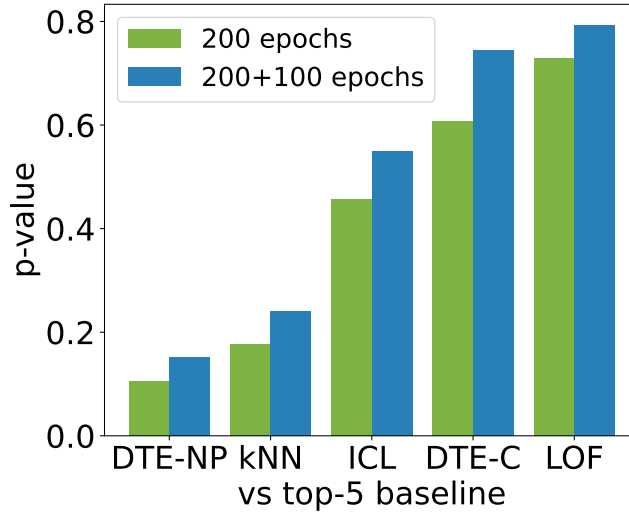


Figure 14: p -values increase with additional 100 epochs of pretraining, i.e. FoMo-0D w/ $D = 100$ performs better against top-5 baselines on datasets w/ $d \leq 100$.

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.

Besides the ablation studies, we provide a qualitative case study of sample-to-sample attention in Appendix H, showing that an outlier attends to the points in context that are within a short distance significantly more than random points, suggesting that PFNs tend to mimic non-parametrics.

K. Generalization Analyses

K.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

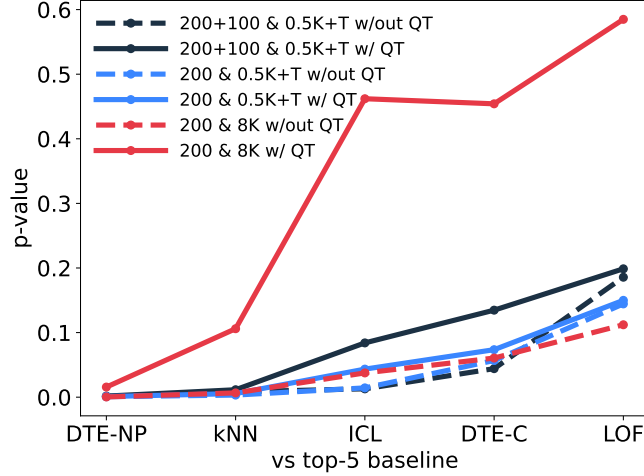


Figure 15: p -values increase, i.e. FoMo-OD performance improves, against top-5 baselines with quantile transform (QT) preceding inference, for 3 different settings of FoMo-OD w/ $D = 100$.

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-OD with a context length of 5000. We present the results with averaged performance over 1000 datasets for each setting in Table 8.

Table 8: Average metric score $\pm$ standard dev. over 1000 seeds for different out-of-distribution (OOD) synthetic GMMs. FoMo-OD 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).

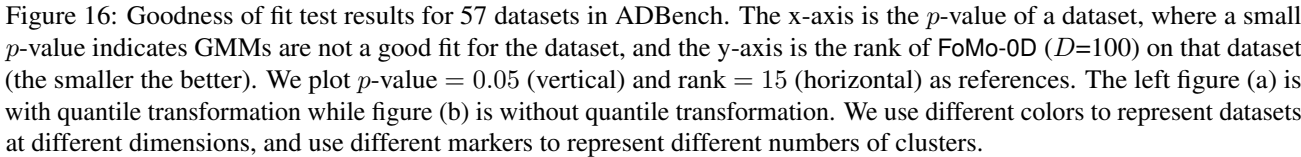
Dataset	AUROC	AUCPR	F1
ID: in-distribution	98.55 $\pm$ 2.73	91.17 $\pm$ 13.07	86.74 $\pm$ 15.43
(a) OOD w.r.t. $ \mu , \Sigma \in [5, 50]$	94.79 $\pm$ 7.53	80.62 $\pm$ 21.19	76.32 $\pm$ 19.85
(b) OOD w.r.t. $m \in [5, 50]$	97.69 $\pm$ 3.59	86.72 $\pm$ 15.57	81.20 $\pm$ 16.23
(c) OOD w.r.t. $d \in [100, 500]$	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-OD 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-OD is only trained on continuous data, it can almost perfectly classify binary outliers hidden in one of the sub-dimensions, suggesting FoMo-OD could potentially generalize to discrete data at test time.

However, with all variations added, FoMo-OD 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.

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

To understand how FoMo-OD 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



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.

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

⁸We use e-test from <https://www.rdocumentation.org/packages/energy/versions/1.7-11/topics/eqdist.etest>

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-OD 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 9 and 10, respectively.

Table 9: 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-OD
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

Table 10: 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-OD
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

We further report the p -value of the Wilcoxon sign test between the baselines and FoMo-OD on the 10 datasets from OpenOOD, as well as on the expanded benchmark combining those 10 with our original ADBench (10+57) in Table 11. In terms of metric values, FoMo-OD performs 2nd or 3rd best across OOD datasets. p -values show that it significantly outperforms DTE-NP and k NN ($p > 0.95$, such that the p -value < 0.05 for rejecting the null hypothesis and accepting the alternative hypothesis that the “baseline-minus-FoMo-OD” gap is smaller than zero) and is no different from ICL (2nd best after DTE-C). These results demonstrate that FoMo-OD generalizes beyond OD datasets and maintains strong zero-shot OD performance on complex, ImageNet-level OOD benchmarks.

Table 11: p -value of the Wilcoxon sign test (alternative: “greater”) between baselines and FoMo-OD 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-OD.

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 k NN on the OpenOOD datasets, whereas on ADBench, DTE-NP and k NN 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 k NN to estimate meaningful decision boundaries in high dimensions (e.g., 512). In contrast, the performance of FoMo-OD is consistently competitive, where the p -values on the combined testbed (OpenOOD+ADBench) show that FoMo-OD is as competitive as all the top baselines across 67 diverse datasets, while maintaining zero-shot detection ability.

L. Performance Profile Plots

To enable a comprehensive comparison of different methods, we 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 17, Figure 18, and Figure 19 illustrate performance profile plots of FoMo-OD and other baselines across all datasets. The results show that FoMo-OD (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-OD (D=100) is even better (i.e., ranked within top-2) when tested on datasets with dimensions less than 100. As shown in Figure 20, Figure 21, and Figure 22, the area under the curve of FoMo-OD (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.

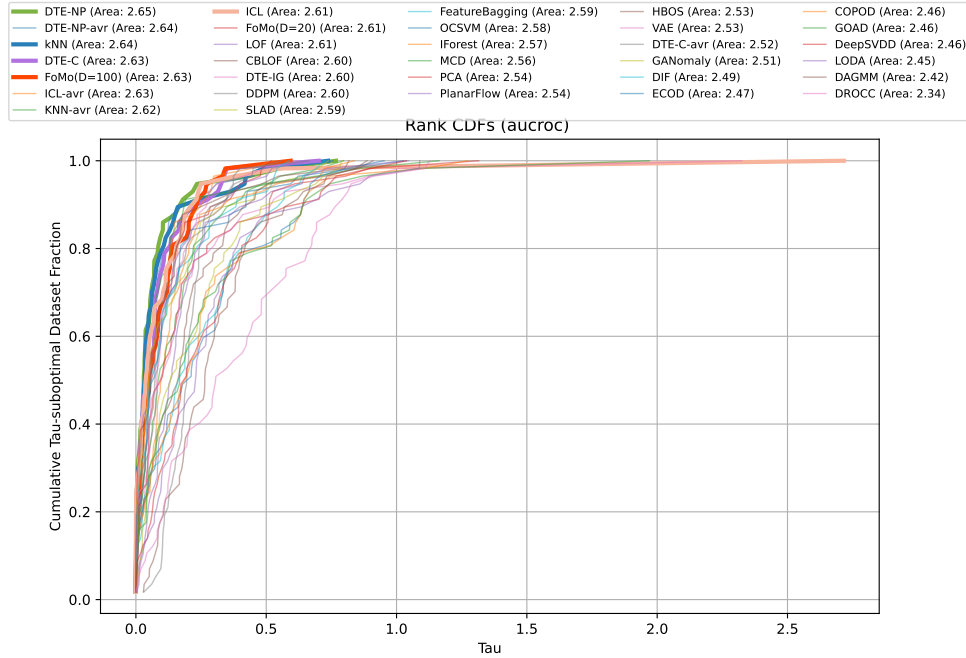


Figure 17: FoMo-OD ranks in **top-5** based on the performance profile plots of all detectors w.r.t. **AUROC** across **all datasets**. In the plot, x-axis represents the τ values—performance ratios that compare each method’s metric to the best performance achieved, while y-axis displays the cumulative 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.

M. Full Results

Tables 13.1 & 13.2, 14.1 & 14.2, and 15.1 & 15.2 respectively show the AUROC, AUPR, and F1 scores of the top-4 baselines, DTE-NP, k NN, ICL, and DTE-C as well as their corresponding avg^{g} model with the average performance across HPs, as listed in Table 2.

Tables 16.1 & 16.2, 17.1 & 17.2, and 18.1 & 18.2 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-OD w/ $D = 100$ with other baselines. The preceding four rows are the same for FoMo-OD w/ $D = 20$, when ranking 31 models (26 baselines + 4 avg^{g} variants of top-4 baselines + FoMo-OD w/ $D = 20$).

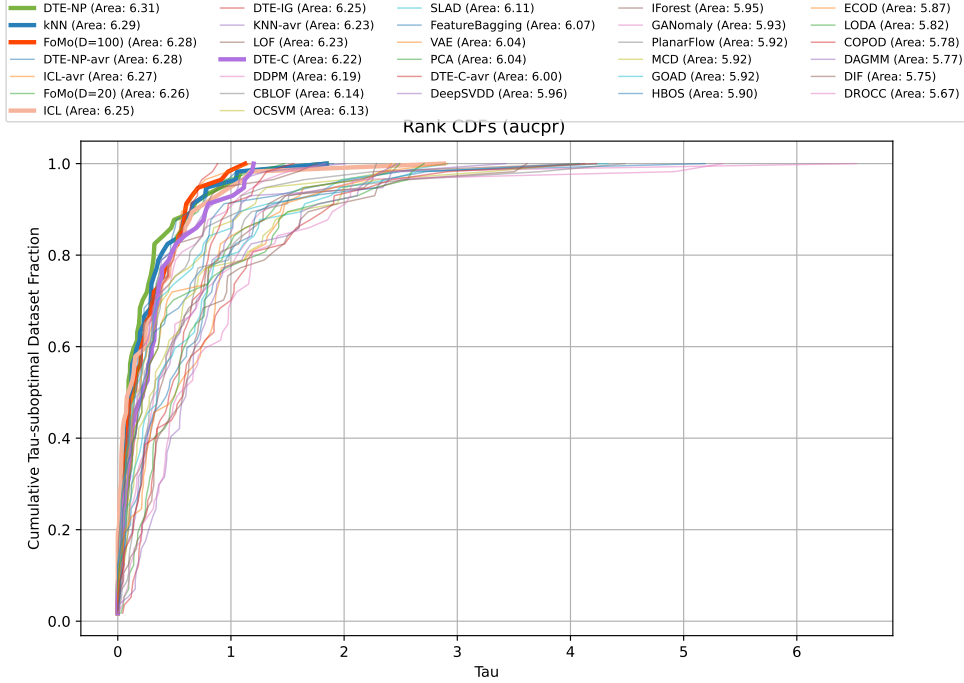


Figure 18: FoMo-0D ranks at **top-3** based on the performance profile plots of all detectors w.r.t. **AUPR** across **all datasets**. In the plot, x-axis represents the τ values—performance ratios that compare each method’s metric to the best performance achieved, while y-axis displays the cumulative 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.

Table 12: 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**³ 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** (kNN , 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	kNN	ICL	DTE-C	LOF	CBLOF	Feat.Bag.	SLAD	DDPM	OCSVM	DTE-NP ^{avg}	kNN ^{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

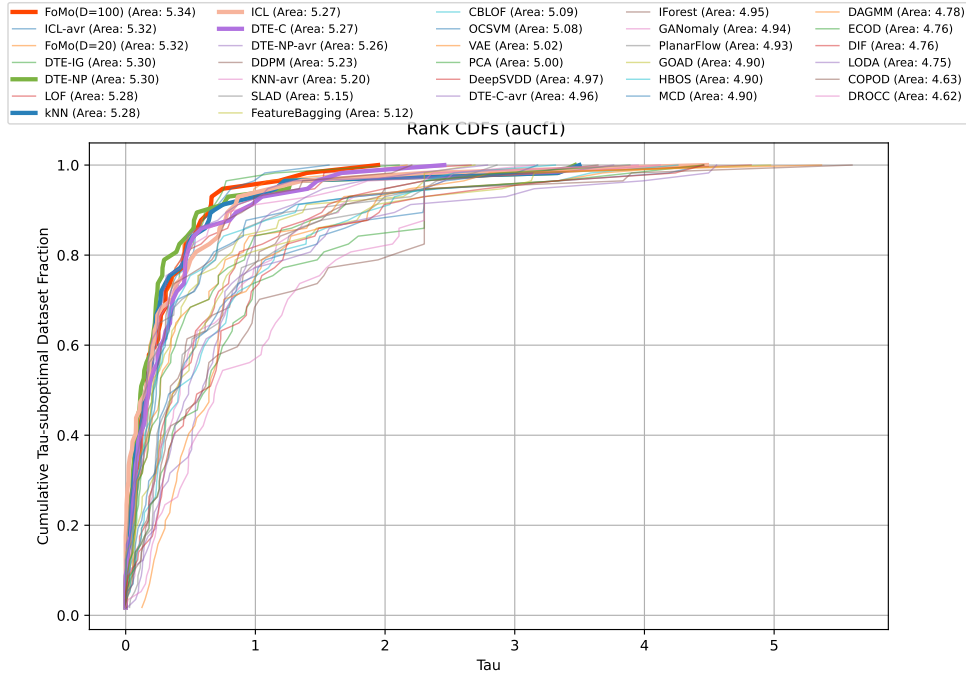


Figure 19: FoMo-0D ranks at **top-1** based on the performance profile plots of all detectors w.r.t. F1 across all datasets. In the plot, x-axis represents the τ values—performance ratios that compare each method’s metric to the best performance achieved, while y-axis displays the cumulative fraction 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.

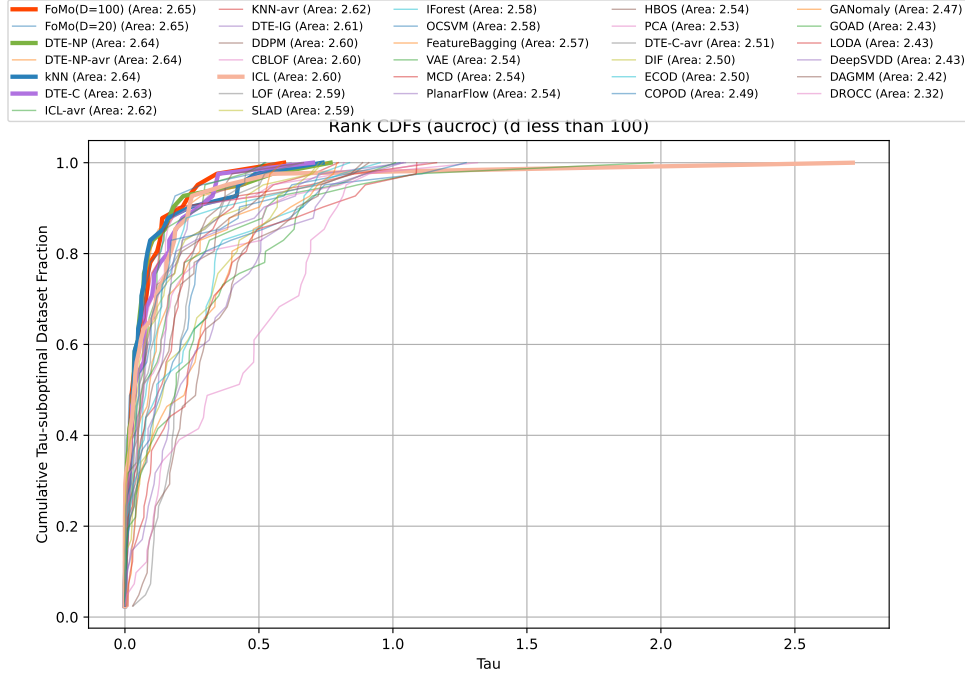


Figure 20: FoMo-OD (D=100) ranks at **top-1** based on the performance profile plots of all detectors w.r.t. **AUROC** in datasets with dimensions less than $d \leq 100$. In the plot, x-axis represents the τ values—performance ratios that compare each method’s metric to the best performance achieved, 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.

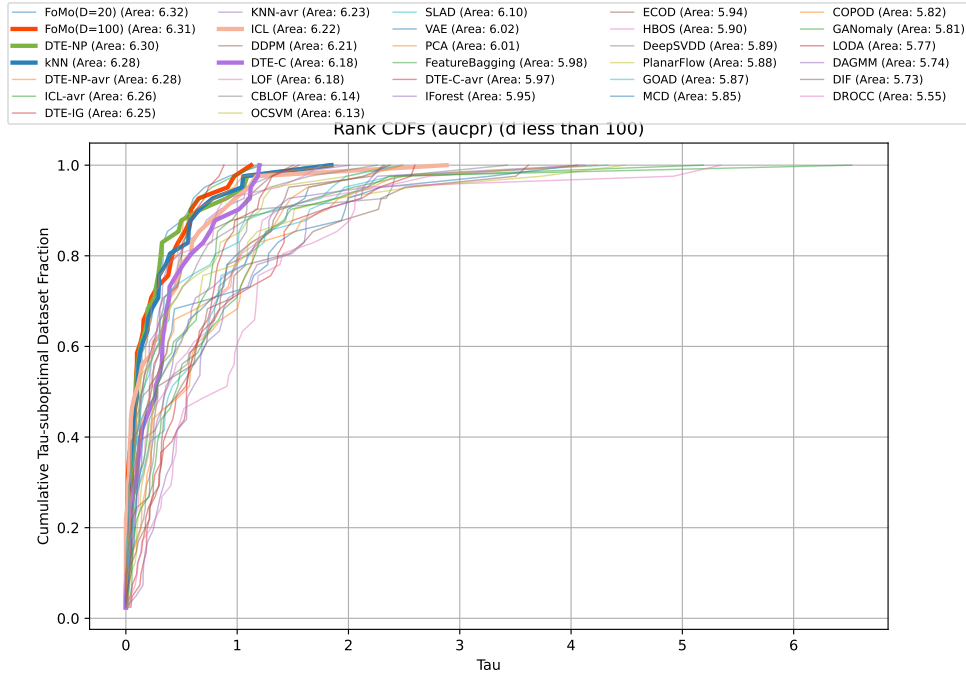


Figure 21: FoMo-0D (D=100) ranks at **top-2** based on the performance profile plots of all detectors w.r.t. **AUPR** in *datasets with dimensions less than $d \leq 100$* . In the plot, x-axis represents the τ values—performance ratios that compare each method’s metric to the best performance achieved, 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.

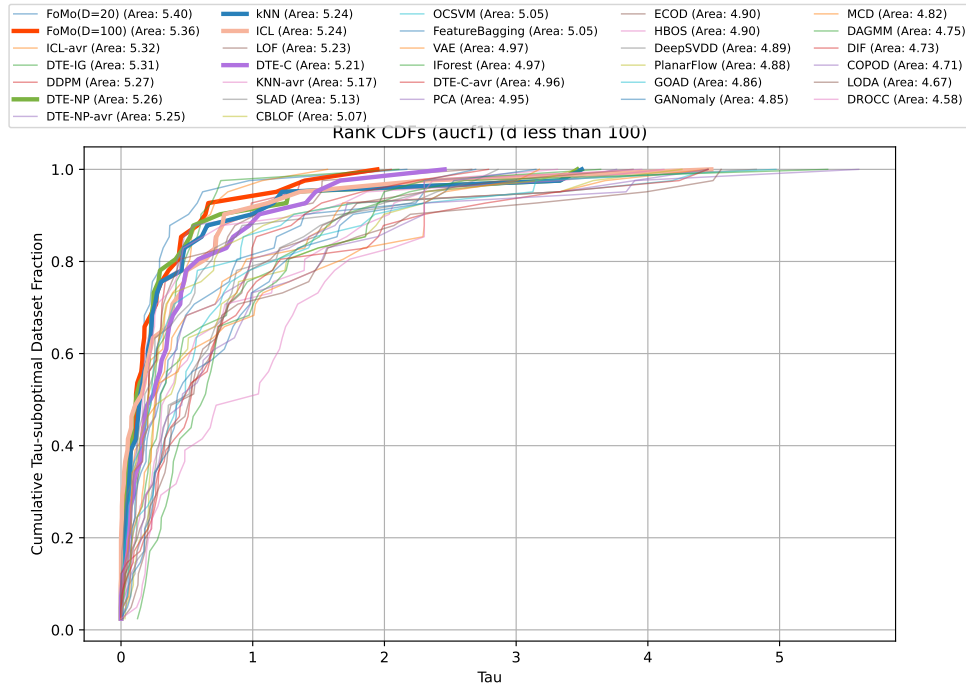


Figure 22: FoMo-0D (D=100) ranks at **top-2** based on the performance profile plots of all detectors w.r.t. **F1** in *datasets with dimensions less than $d \leq 100$* . In the plot, x-axis represents the τ values—performance ratios that compare each method’s metric to the best performance achieved, 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.

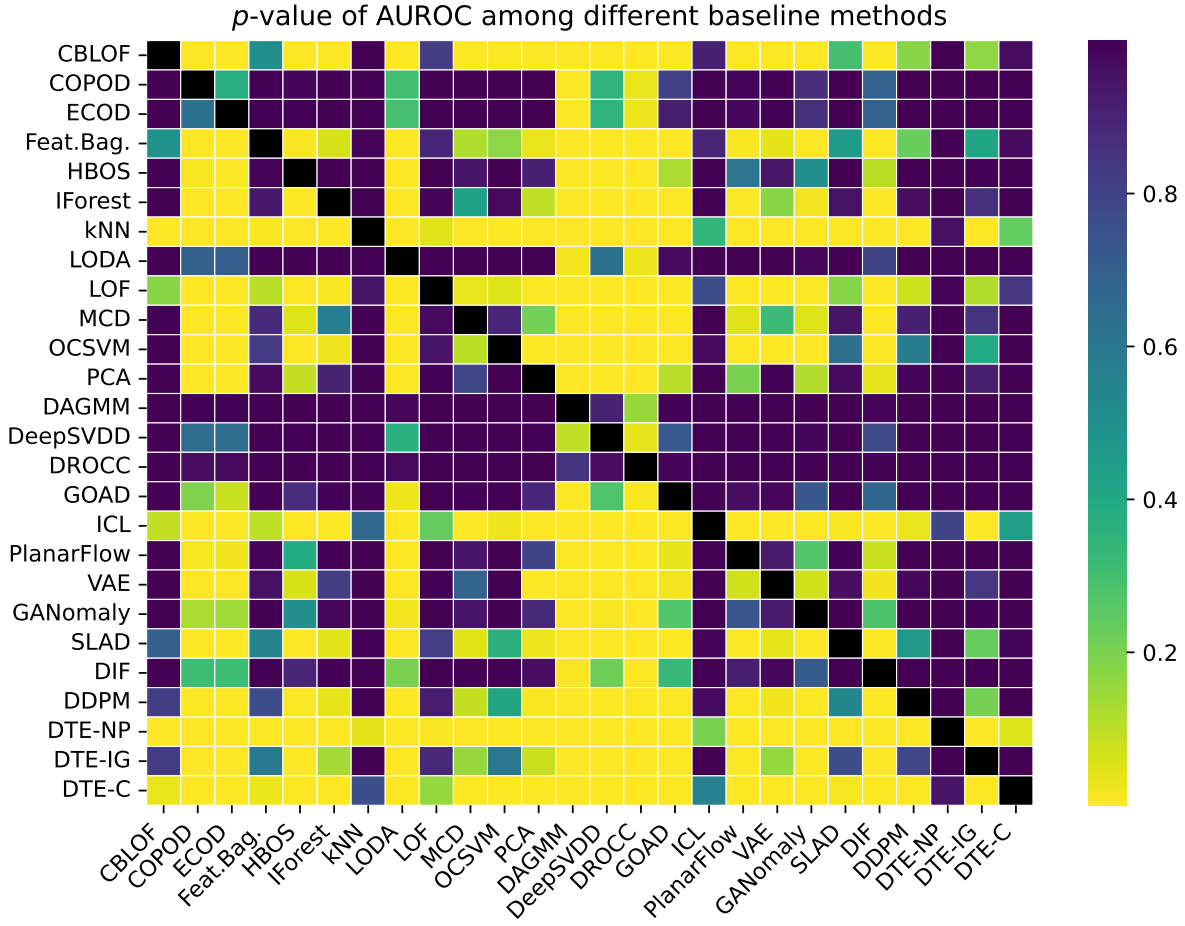


Figure 23: Pairwise *p*-values among baseline methods based on the Wilcoxon signed rank test w.r.t. AUROC performances across datasets.

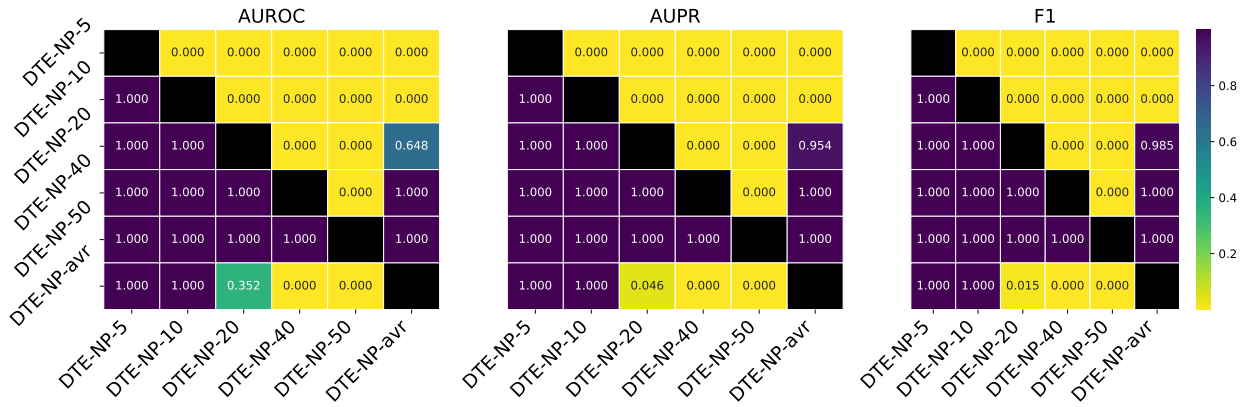


Figure 24: *p*-values w.r.t. AUROC/AUPR/F1 among different HP configurations of **DTE-NP** (i.e., $k \in \{5, 10, 20, 40, 50\}$), along with the ^{avg} model with the average performance across HPs.

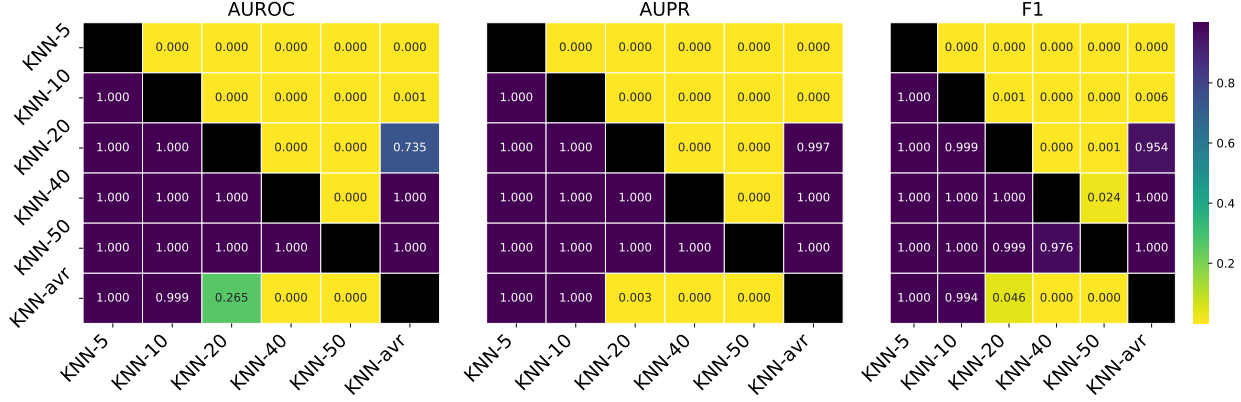


Figure 25: p -values w.r.t. AUROC/AUPR/F1 among different HP configurations of k NN (i.e., $k \in \{5, 10, 20, 40, 50\}$), along with the avg model with the average performance across HPs.

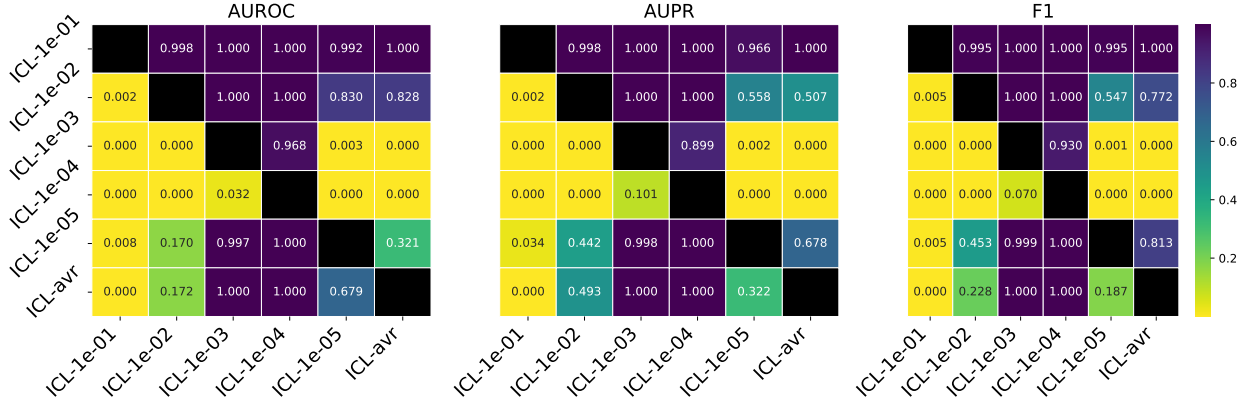


Figure 26: p -values w.r.t. AUROC/AUPR/F1 among different HP configurations of ICL (i.e., learning_rate $\in \{10^{-1}, 10^{-2}, 10^{-3}, 10^{-4}, 10^{-5}\}$), along with the avg model with the average performance across HPs.

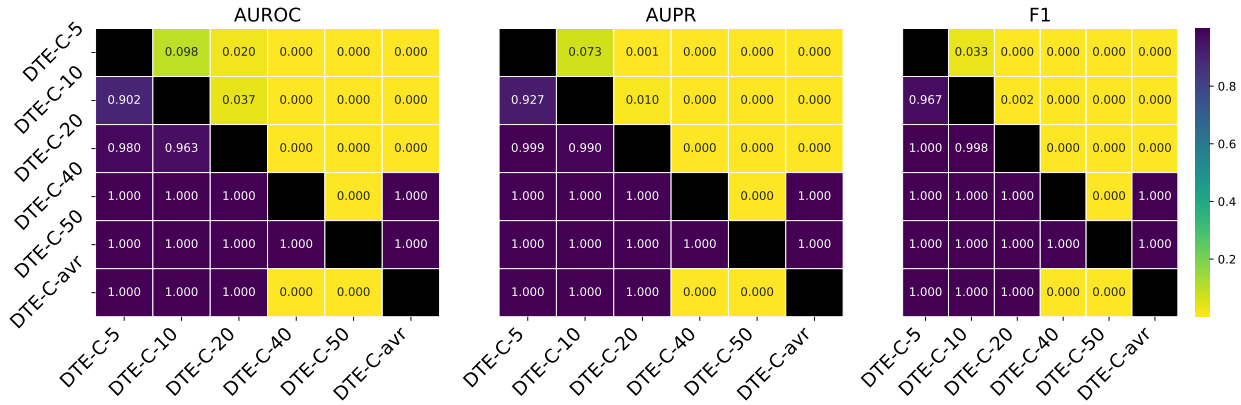


Figure 27: p -values w.r.t. AUROC/AUPR/F1 among different HP configurations of DTE-C (i.e., $k \in \{5, 10, 20, 40, 50\}$), along with the avg model with the average performance across HPs.

Table 13.1: 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
aloi	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	60.52 $\pm$ 0.00	60.58 $\pm$ 0.00	60.52 $\pm$ 0.00	60.52 $\pm$ 0.00	60.02 $\pm$ 0.00	60.23 $\pm$ 0.00	60.25 $\pm$ 0.00
amthroid	93.01 $\pm$ 0.00	92.89 $\pm$ 0.00	92.66 $\pm$ 0.00	92.38 $\pm$ 0.00	92.26 $\pm$ 0.00	92.64 $\pm$ 0.00	92.81 $\pm$ 0.00	92.60 $\pm$ 0.00	92.34 $\pm$ 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	90.68 $\pm$ 0.45	92.55 $\pm$ 0.44	93.71 $\pm$ 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
breasw	99.10 $\pm$ 0.28	98.91 $\pm$ 0.34	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	78.60 $\pm$ 0.00	78.66 $\pm$ 0.00
cardio	91.53 $\pm$ 0.00	91.03 $\pm$ 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	93.85 $\pm$ 0.00	94.08 $\pm$ 0.00	93.07 $\pm$ 0.00
cardiography	60.40 $\pm$ 0.00	61.63 $\pm$ 0.00	63.14 $\pm$ 0.00	64.64 $\pm$ 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	72.58 $\pm$ 0.26	74.81 $\pm$ 0.34	76.87 $\pm$ 0.38	77.47 $\pm$ 0.37	74.42 $\pm$ 0.36	72.91 $\pm$ 0.29	75.24 $\pm$ 0.40	77.50 $\pm$ 0.47	79.14 $\pm$ 0.37	79.68 $\pm$ 0.37	76.90 $\pm$ 0.35
cinema	70.39 $\pm$ 0.33	72.34 $\pm$ 0.17	72.28 $\pm$ 0.10	71.93 $\pm$ 0.17	71.80 $\pm$ 0.16	72.11 $\pm$ 0.16	72.23 $\pm$ 0.15	72.36 $\pm$ 0.12	71.94 $\pm$ 0.19	71.37 $\pm$ 0.21	71.28 $\pm$ 0.16	71.84 $\pm$ 0.15
cover	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 $\pm$ 0.31	96.73 $\pm$ 0.22
donors	97.72 $\pm$ 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	97.51 $\pm$ 0.06	99.24 $\pm$ 0.08	98.85 $\pm$ 0.10	98.20 $\pm$ 0.13	97.90 $\pm$ 0.14	98.74 $\pm$ 0.09
fault	58.32 $\pm$ 0.00	58.37 $\pm$ 0.00	58.70 $\pm$ 0.00	60.00 $\pm$ 0.00	60.45 $\pm$ 0.00	59.17 $\pm$ 0.00	58.23 $\pm$ 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	96.70 $\pm$ 0.90	95.67 $\pm$ 0.93	95.64 $\pm$ 0.93	95.60 $\pm$ 0.92	95.60 $\pm$ 0.92	95.64 $\pm$ 0.92	95.59 $\pm$ 0.97	95.55 $\pm$ 0.98	95.53 $\pm$ 0.89	95.54 $\pm$ 0.92	95.62 $\pm$ 0.88	95.57 $\pm$ 0.93
glass	96.08 $\pm$ 0.39	95.04 $\pm$ 1.06	95.82 $\pm$ 1.12	97.89 $\pm$ 1.12	87.31 $\pm$ 1.40	90.83 $\pm$ 0.91	92.13 $\pm$ 0.94	88.67 $\pm$ 0.98	87.24 $\pm$ 1.18	84.93 $\pm$ 2.92	83.55 $\pm$ 2.61	87.30 $\pm$ 1.59
hepatitis	99.84 $\pm$ 0.20	99.27 $\pm$ 0.51	96.89 $\pm$ 0.96	93.15 $\pm$ 1.69	91.97 $\pm$ 1.76	96.22 $\pm$ 0.88	96.77 $\pm$ 1.47	86.88 $\pm$ 2.21	85.50 $\pm$ 2.34	85.46 $\pm$ 1.92	84.88 $\pm$ 2.09	87.90 $\pm$ 1.75
hup	99.99 $\pm$ 0.00	99.98 $\pm$ 0.01	99.95 $\pm$ 0.01	99.93 $\pm$ 0.01	99.92 $\pm$ 0.02	99.95 $\pm$ 0.01	100.00 $\pm$ 0.00	99.99 $\pm$ 0.02	99.95 $\pm$ 0.01	99.95 $\pm$ 0.01	99.96 $\pm$ 0.01	99.96 $\pm$ 0.01
imdb	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	50.35 $\pm$ 0.00	68.08 $\pm$ 0.00	50.04 $\pm$ 0.00	50.29 $\pm$ 0.00	50.23 $\pm$ 0.00	50.23 $\pm$ 0.00	50.18 $\pm$ 0.00
internets	70.96 $\pm$ 0.60	68.65 $\pm$ 0.00	66.86 $\pm$ 0.00	65.97 $\pm$ 0.00	65.83 $\pm$ 0.00	67.65 $\pm$ 0.00	68.08 $\pm$ 0.00	65.48 $\pm$ 0.00	65.02 $\pm$ 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 $\pm$ 0.74	97.84 $\pm$ 0.64	96.83 $\pm$ 0.71	96.21 $\pm$ 0.79	97.50 $\pm$ 0.63	97.62 $\pm$ 0.81	97.32 $\pm$ 0.85	96.33 $\pm$ 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	64.36 $\pm$ 0.00	62.49 $\pm$ 0.00	61.93 $\pm$ 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	33.40 $\pm$ 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	99.79 $\pm$ 0.32	99.79 $\pm$ 0.32	99.76 $\pm$ 0.31	99.76 $\pm$ 0.31	99.80 $\pm$ 0.30	93.27 $\pm$ 0.00	99.87 $\pm$ 0.10	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	83.49 $\pm$ 0.00	82.87 $\pm$ 0.00	82.05 $\pm$ 0.00	81.73 $\pm$ 0.00	82.81 $\pm$ 0.00	83.27 $\pm$ 0.00	82.64 $\pm$ 0.00	81.83 $\pm$ 0.00	80.76 $\pm$ 0.00	80.30 $\pm$ 0.00	81.76 $\pm$ 0.00
mammography	87.65 $\pm$ 0.00	87.73 $\pm$ 0.00	87.68 $\pm$ 0.00	87.42 $\pm$ 0.00	87.29 $\pm$ 0.00	87.55 $\pm$ 0.00	93.85 $\pm$ 0.00	87.75 $\pm$ 0.00	87.38 $\pm$ 0.00	86.97 $\pm$ 0.00	86.78 $\pm$ 0.00	87.29 $\pm$ 0.00
minst	94.22 $\pm$ 0.00	93.93 $\pm$ 0.00	93.57 $\pm$ 0.00	93.20 $\pm$ 0.00	93.08 $\pm$ 0.00	93.60 $\pm$ 0.00	93.85 $\pm$ 0.00	93.45 $\pm$ 0.00	93.00 $\pm$ 0.00	92.55 $\pm$ 0.00	92.36 $\pm$ 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
optdigits	95.00 $\pm$ 0.00	93.97 $\pm$ 0.00	92.52 $\pm$ 0.00	90.87 $\pm$ 0.00	90.28 $\pm$ 0.00	92.53 $\pm$ 0.00	93.72 $\pm$ 0.00	92.09 $\pm$ 0.00	90.20 $\pm$ 0.00	87.66 $\pm$ 0.00	86.62 $\pm$ 0.00	90.07 $\pm$ 0.00
pagesblocks	89.04 $\pm$ 0.00	89.40 $\pm$ 0.00	89.56 $\pm$ 0.00	89.37 $\pm$ 0.00	89.23 $\pm$ 0.00	89.32 $\pm$ 0.00	89.65 $\pm$ 0.00	89.86 $\pm$ 0.00	89.88 $\pm$ 0.00	89.30 $\pm$ 0.00	89.18 $\pm$ 0.00	89.57 $\pm$ 0.00
pendigits	99.90 $\pm$ 0.00	99.88 $\pm$ 0.00	99.83 $\pm$ 0.00	99.51 $\pm$ 0.00	99.38 $\pm$ 0.00	99.70 $\pm$ 0.00	99.87 $\pm$ 0.00	99.79 $\pm$ 0.00	99.21 $\pm$ 0.00	98.58 $\pm$ 0.00	98.39 $\pm$ 0.00	99.17 $\pm$ 0.00
pinas	82.21 $\pm$ 1.82	79.74 $\pm$ 1.61	77.98 $\pm$ 1.38	77.28 $\pm$ 1.35	77.14 $\pm$ 1.33	78.87 $\pm$ 1.43	77.44 $\pm$ 2.07	76.14 $\pm$ 1.36	76.39 $\pm$ 1.21	76.55 $\pm$ 1.25	76.38 $\pm$ 1.29	76.58 $\pm$ 1.39
satellite	82.40 $\pm$ 0.00	82.09 $\pm$ 0.00	81.55 $\pm$ 0.00	80.71 $\pm$ 0.00	80.38 $\pm$ 0.00	81.43 $\pm$ 0.00	82.24 $\pm$ 0.00	81.56 $\pm$ 0.00	80.56 $\pm$ 0.00	79.22 $\pm$ 0.00	78.76 $\pm$ 0.00	80.47 $\pm$ 0.00
satimage-2	92.68 $\pm$ 0.00	92.73 $\pm$ 0.00	92.79 $\pm$ 0.00	92.79 $\pm$ 0.00	92.82 $\pm$ 0.00	92.77 $\pm$ 0.00	92.71 $\pm$ 0.00	92.80 $\pm$ 0.00	92.79 $\pm$ 0.00	92.78 $\pm$ 0.00	92.77 $\pm$ 0.00	92.77 $\pm$ 0.00
shuttle	99.94 $\pm$ 0.00	99.92 $\pm$ 0.00	99.92 $\pm$ 0.00	99.91 $\pm$ 0.00	99.91 $\pm$ 0.00	99.92 $\pm$ 0.00	99.91 $\pm$ 0.00	99.91 $\pm$ 0.00	99.89 $\pm$ 0.00	99.89 $\pm$ 0.00	99.89 $\pm$ 0.00	99.89 $\pm$ 0.00
skin	99.66 $\pm$ 0.05	99.52 $\pm$ 0.04	99.28 $\pm$ 0.02	98.77 $\pm$ 0.09	98.53 $\pm$ 0.09	99.15 $\pm$ 0.05	99.51 $\pm$ 0.06	99.17 $\pm$ 0.05	98.70 $\pm$ 0.12	97.43 $\pm$ 0.05	96.90 $\pm$ 0.09	98.34 $\pm$ 0.04
snip	92.94 $\pm$ 2.55	92.89 $\pm$ 2.56	93.14 $\pm$ 2.20	93.14 $\pm$ 2.15	93.14 $\pm$ 2.20	93.04 $\pm$ 2.32	92.90 $\pm$ 2.48	92.97 $\pm$ 2.52	93.25 $\pm$ 2.06	93.11 $\pm$ 2.28	93.24 $\pm$ 2.25	93.10 $\pm$ 2.30
spambase	84.06 $\pm$ 0.00	83.49 $\pm$ 0.00	82.97 $\pm$ 0.00	82.50 $\pm$ 0.00	82.36 $\pm$ 0.00	83.07 $\pm$ 0.00	83.36 $\pm$ 0.00	82.68 $\pm$ 0.00	82.22 $\pm$ 0.00	81.83 $\pm$ 0.00	81.72 $\pm$ 0.00	82.36 $\pm$ 0.00
speech	41.12 $\pm$ 0.00	39.00 $\pm$ 0.00	37.59 $\pm$ 0.00	37.37 $\pm$ 0.00	37.01 $\pm$ 0.00	38.42 $\pm$ 0.00	36.36 $\pm$ 0.00	36.15 $\pm$ 0.00	36.37 $\pm$ 0.00	36.40 $\pm$ 0.00	36.25 $\pm$ 0.00	36.31 $\pm$ 0.00
stamps	97.88 $\pm$ 0.33	97.04 $\pm$ 0.33	95.60 $\pm$ 0.78	94.59 $\pm$ 1.07	94.29 $\pm$ 1.08	95.88 $\pm$ 0.66	96.19 $\pm$ 1.24	94.35 $\pm$ 1.35	93.67 $\pm$ 1.09	93.44 $\pm$ 1.21	93.33 $\pm$ 1.25	94.20 $\pm$ 1.17
thyroid	98.58 $\pm$ 0.00	98.63 $\pm$ 0.00	98.63 $\pm$ 0.00	98.63 $\pm$ 0.00	98.59 $\pm$ 0.00	98.61 $\pm$ 0.00	98.68 $\pm$ 0.00	98.67 $\pm$ 0.00	98.69 $\pm$ 0.00	98.69 $\pm$ 0.00	98.69 $\pm$ 0.00	98.68 $\pm$ 0.00
verbal	97.59 $\pm$ 2.23	67.05 $\pm$ 2.09	56.99 $\pm$ 2.07	49.07 $\pm$ 1.60	47.08 $\pm$ 1.43	59.96 $\pm$ 1.83	57.33 $\pm$ 3.80	49.40 $\pm$ 1.87	45.06 $\pm$ 1.20	41.08 $\pm$ 1.52	40.08 $\pm$ 1.23	46.59 $\pm$ 1.67
vowels	82.11 $\pm$ 0.00	81.62 $\pm$ 0.00	80.46 $\pm$ 0.00	78.11 $\pm$ 0.00	76.87 $\pm$ 0.00	79.83 $\pm$ 0.00	82.21 $\pm$ 0.00	80.20 $\pm$ 0.00	77.82 $\pm$ 0.00	72.26 $\pm$ 0.00	69.03 $\pm$ 0.00	76.30 $\pm$ 0.00
waveform	74.42 $\pm$ 0.00	75.40 $\pm$ 0.00	76.12 $\pm$ 0.00	76.79 $\pm$ 0.00	76.90 $\pm$ 0.00	75.93 $\pm$ 0.00	75.21 $\pm$ 0.00	76.04 $\pm$ 0.00	76.78 $\pm$ 0.00	77.47 $\pm$ 0.00	77.60 $\pm$ 0.00	76.62 $\pm$ 0.00
wbc	99.62 $\pm$ 0.27	99.36 $\pm$ 0.33	99.17 $\pm$ 0.41	99.12 $\pm$ 0.37	99.09 $\pm$ 0.38	99.27 $\pm$ 0.34	99.23 $\pm$ 0.27	99.09 $\pm$ 0.29	99.08 $\pm$ 0.21	99.08		

Table 13.2: 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 respectively to mark the **best** and the **worst** performance of each model to showcase the variability across different HP settings.

dataset	ICL-0.1	ICL-0.01	ICL-0.001	ICL-0.0001	ICL-1e-05	ICL-avr	DTE-C-5	DTE-C-10	DTE-C-20	DTE-C-40	DTE-C-50	DTE-C-avg
aloi	47.74±0.28	47.12±0.51	46.81±0.47	48.42±0.24	48.06±0.23	47.63±0.15	50.20±0.21	50.84±0.20	50.16±0.34	50.26±0.33	50.00±0.00	50.29±0.09
amazon	53.07±0.19	53.44±0.19	52.75±0.68	53.31±0.21	53.18±0.15	53.15±0.19	56.20±1.62	55.07±0.20	56.12±0.73	50.00±0.00	50.00±0.00	53.48±1.33
amthroid	84.02±0.46	72.68±3.79	87.29±2.69	88.84±2.35	88.52±1.40	84.27±2.31	92.64±0.15	97.47±0.10	92.64±0.15	97.40±0.22	50.00±0.00	88.05±0.05
backdoor	93.03±0.66	92.91±0.70	93.30±0.44	93.93±0.58	93.32±0.41	93.30±0.48	88.65±1.08	92.06±0.14	92.64±0.15	98.85±0.45	99.26±0.07	74.67±0.36
breastw	98.96±0.47	98.87±0.20	99.19±0.34	99.11±0.18	97.61±0.38	98.75±0.18	93.45±1.34	94.34±1.25	98.85±0.45	98.85±0.45	99.26±0.07	96.52±0.53
campaign	76.07±0.87	74.61±1.97	78.88±0.42	79.79±0.91	82.30±0.52	78.33±0.52	79.18±1.09	78.24±0.99	78.49±1.13	50.00±0.00	50.00±0.00	67.18±0.74
cardio	73.99±7.79	84.98±4.75	82.34±3.36	78.68±2.18	68.59±0.64	77.71±1.76	88.07±0.51	87.54±0.69	87.66±0.63	50.00±0.00	50.00±0.00	72.66±0.21
cardiotocography	48.99±2.02	54.18±2.77	53.24±3.25	50.67±2.55	47.20±1.33	50.85±1.24	60.36±1.54	60.09±2.19	59.08±1.34	50.00±0.00	50.00±0.00	55.90±0.23
celeba	79.38±2.17	79.15±2.47	76.13±1.88	78.39±1.78	79.43±1.44	78.49±0.93	82.95±0.91	81.59±0.79	80.16±1.26	50.00±0.00	50.00±0.00	68.94±0.44
census	70.41±2.07	67.52±8.79	73.93±0.78	75.03±0.46	74.37±0.53	72.25±1.63	70.95±0.91	68.04±0.85	68.04±0.85	50.00±0.00	50.00±0.00	61.41±0.71
cover	93.59±3.09	82.32±9.14	91.92±4.58	99.37±0.30	99.51±0.14	91.42±1.74	97.57±0.86	97.56±0.53	95.64±0.27	58.94±17.89	50.00±0.00	78.40±0.14
donors	86.20±10.29	97.76±1.57	98.64±0.55	99.37±0.30	99.51±0.14	96.30±1.19	98.68±0.15	97.56±0.53	95.64±0.27	58.94±17.89	50.00±0.00	80.17±3.55
fault	63.37±3.29	63.04±1.93	61.65±0.61	61.44±0.87	62.83±1.01	62.47±1.14	94.19±1.69	58.41±1.41	59.04±1.53	50.21±0.25	50.21±0.25	55.41±0.63
fraud	93.24±1.45	93.21±1.25	94.97±0.71	95.84±0.87	95.90±1.05	94.22±1.45	94.49±1.56	93.08±2.20	93.50±2.01	50.00±0.00	50.00±0.00	76.21±1.08
glass	84.02±8.38	94.06±5.83	99.25±0.37	99.30±0.35	99.12±0.55	95.27±1.45	93.46±0.91	89.96±2.70	84.89±2.89	66.42±8.60	50.00±0.00	76.95±2.58
hepatitis	99.95±0.11	99.97±0.02	99.93±0.14	99.86±0.29	99.97±0.06	99.98±0.02	99.54±0.58	98.90±1.01	98.90±1.08	99.39±0.07	50.00±0.00	89.59±0.09
hup	99.96±0.07	99.97±0.01	99.98±0.01	100.00±0.00	99.97±0.06	99.98±0.02	99.54±0.58	98.90±1.01	98.90±1.08	99.39±0.07	50.00±0.00	89.59±0.09
indb	52.16±0.31	51.88±0.40	52.28±0.15	52.35±0.12	53.06±0.15	52.35±0.12	47.68±3.79	50.89±1.90	47.81±2.08	50.00±0.00	50.00±0.00	49.27±1.12
ionosphere	96.81±2.22	96.13±3.45	98.90±0.41	98.91±0.31	98.14±0.79	97.78±1.23	94.67±1.52	94.49±2.41	95.23±1.37	85.77±2.62	44.90±23.57	83.01±5.81
landsat	65.71±2.05	60.63±1.79	65.95±0.76	67.85±0.68	61.82±1.06	64.39±0.59	51.80±0.94	37.72±1.31	48.62±2.62	50.22±4.45	50.00±0.00	42.38±0.25
letter	48.14±3.53	42.74±3.41	40.70±2.13	40.32±1.11	47.20±2.50	43.82±1.97	36.75±1.19	37.72±1.31	37.72±1.31	50.00±0.00	50.00±0.00	88.19±1.78
lymphography	100.00±0.00	100.00±0.00	100.00±0.00	100.00±0.00	100.00±0.00	100.00±0.00	98.97±0.31	99.40±0.31	98.42±0.59	93.82±8.49	50.00±0.00	75.19±3.50
magic gamma	69.73±3.02	76.94±3.25	77.61±0.70	77.81±0.88	78.36±1.34	80.47±0.62	86.42±0.59	87.42±0.59	84.88±18.23	83.40±1.03	50.00±0.00	88.19±1.78
mnist	79.90±5.19	85.08±3.05	79.26±2.47	80.75±1.40	77.35±0.70	76.09±1.17	86.42±0.59	87.42±0.59	84.88±18.23	83.40±1.03	50.00±0.00	75.19±3.50
mnist	75.53±1.43	76.05±3.81	79.13±1.20	87.05±1.12	89.10±0.77	81.37±0.93	90.21±0.70	86.98±1.70	85.58±2.12	50.00±0.00	50.00±0.00	72.55±0.64
mnist	100.00±0.00	100.00±0.00	100.00±0.00	100.00±0.00	100.00±0.00	100.00±0.00	90.21±0.70	86.98±1.70	85.58±2.12	50.00±0.00	50.00±0.00	72.55±0.64
mnist	91.22±1.88	94.15±0.97	95.17±0.85	98.47±0.03	95.08±1.06	94.82±0.49	86.29±1.29	74.81±6.84	77.57±4.38	50.00±0.00	50.00±0.00	67.73±1.71
optdigits	89.71±1.10	90.77±1.73	88.70±1.01	88.58±0.56	95.10±0.89	89.16±0.29	90.29±1.29	89.29±1.29	88.80±0.13	50.00±0.00	50.00±0.00	81.66±0.20
pagesblocks	87.01±9.78	96.45±2.40	96.17±2.68	98.28±0.77	95.10±0.89	94.60±0.36	98.24±0.47	97.42±0.83	96.89±0.68	50.00±0.00	50.00±0.00	78.51±0.15
pin	74.66±3.27	73.04±1.77	79.77±1.92	78.06±0.74	75.09±2.65	76.12±1.04	80.21±1.18	79.01±0.76	78.54±0.62	56.11±12.22	50.00±0.00	68.77±2.26
satellite	74.62±9.67	83.23±2.89	85.43±0.52	89.24±0.57	85.35±0.48	83.57±2.48	99.61±0.13	99.17±0.15	98.38±0.56	69.34±23.69	50.00±0.00	83.30±4.71
satimage-2	94.41±2.46	93.04±12.69	99.80±0.06	99.99±0.00	99.97±0.01	99.85±0.11	99.76±0.01	99.74±0.01	99.70±0.01	99.28±0.22	50.00±0.00	89.70±0.05
shuttle	99.43±0.50	99.89±0.05	99.99±0.00	99.99±0.00	99.97±0.01	99.85±0.11	99.76±0.01	99.74±0.01	99.70±0.01	99.28±0.22	50.00±0.00	89.70±0.05
skin	53.71±32.40	86.42±5.89	86.94±5.22	71.25±17.43	70.72±13.69	88.37±5.41	95.06±1.20	95.61±1.24	95.63±0.28	95.84±1.23	50.00±0.00	83.48±0.19
snip	81.23±10.57	87.97±4.57	90.53±5.25	89.48±4.52	92.63±0.46	80.71±0.86	82.98±0.37	83.29±0.49	83.74±0.38	95.84±1.23	50.00±0.00	86.47±0.96
speech	50.06±3.03	50.96±2.48	53.72±1.46	51.29±2.33	46.64±1.97	50.53±1.11	38.02±1.49	38.55±1.07	38.72±1.79	50.00±0.00	50.00±0.00	43.06±0.55
stamps	77.62±9.15	84.33±6.20	97.29±0.48	96.70±0.85	94.64±3.19	90.11±2.02	93.01±1.38	98.75±0.07	98.92±0.02	90.17±4.76	50.10±26.50	82.21±5.22
thyroid	94.79±1.60	95.24±1.04	96.20±0.74	96.78±1.19	93.73±0.80	95.35±0.60	98.75±0.07	98.92±0.02	98.92±0.02	98.94±0.05	50.00±0.00	59.97±2.61
verbal	53.80±4.97	58.76±7.63	75.60±6.18	82.96±2.86	81.92±1.33	70.61±1.11	67.07±2.73	65.09±3.50	62.99±4.76	56.35±4.78	48.40±5.99	72.28±0.42
vowels	73.59±6.94	79.11±2.66	84.59±3.49	84.81±3.61	83.99±1.93	81.22±1.57	67.07±2.73	65.09±3.50	62.99±4.76	56.35±4.78	48.40±5.99	72.28±0.42
waveform	70.94±1.78	73.85±6.07	70.23±2.82	61.96±1.35	64.69±1.66	68.33±2.15	63.97±2.75	65.16±1.34	63.97±2.75	66.24±1.45	50.00±0.00	59.07±0.68
wbc	96.66±2.87	99.05±0.72	99.51±0.25	99.67±0.15	99.29±0.25	98.84±0.43	84.96±0.46	98.76±0.26	99.01±0.41	50.22±16.16	45.17±9.66	78.42±3.73
wilt	52.75±14.46	61.09±5.94	83.81±3.04	84.63±1.21	79.81±0.93	72.42±2.57	87.06±0.69	81.09±1.72	84.90±1.24	90.99±1.24	77.44±0.66	77.44±0.66
wine	91.76±2.53	90.73±0.29	99.92±0.12	99.81±0.39	99.32±0.33	99.79±0.33	99.91±0.11	99.70±0.33	99.81±0.29	66.04±29.91	50.00±0.00	81.66±0.20
wrbc	91.76±2.53	90.73±0.29	99.92±0.12	99.81±0.39	99.32±0.33	99.79±0.33	99.91±0.11	99.70±0.33	99.81±0.29	66.04±29.91	50.00±0.00	81.66±0.20
wpbc	91.76±2.53	90.73±0.29	99.92±0.12	99.81±0.39	99.32±0.33	99.79±0.33	99.91±0.11	99.70±0.33	99.81±0.29	66.04±29.91	50.00±0.00	81.66±0.20
yeast	47.02±1.77	47.02±1.77	47.25±0.62	47.38±1.07	46.32±0.91	46.62±1.16	46.89±1.26	46.48±1.44	46.48±1.44	44.42±1.06	50.00±0.00	61.79±0.36
yelp	54.26±0.22	54.08±0.78	54.04±0.78	54.07±0.26	52.73±0.21	53.71±0.22	62.00±1.65	59.63±3.66	60.99±2.26	50.00±0.00	50.00±0.00	56.47±0.84
MINST-C	80.78±0.49	83.24±0.51	85.02±0.14	84.91±0.11	83.37±0.12	83.46±0.11	82.43±0.19	86.12±0.43	85.25±0.30	50.00±0.00	50.00±0.00	71.28±0.16
FashionMNIST	80.44±0.09	90.91±0.06	90.91±0.06	90.95±0.03	89.57±0.13	92.90±0.15	90.13±0.24	68.59±0.36	68.87±0.31	50.00±0.00	50.00±0.00	61.25±0.12
CIFAR10	57.83±0.21	64.12±0.37	65.72±0.09	65.89±0.27	59.25±0.13	62.90±0.15	68.78±0.49	69.02±0.33	69.02±0.33	50.00±0.00	50.00±0.00	57.80±0.05
SVHN	59.44±0.27	61.25±0.11	61.73±0.07	61.54±0.06	60.91±0.08	60.91±0.08	63.02±0.30	63.02±0.30	62.97±0.38	50.00±0.00	50.00±0.00	73.32±0.42
MTec-AD	93.31±0.36	93.73±0.30	94.26±0.35	94.27±0.32	89.82±0.44	94.08±0.33	68.78±0.49	69.56±0.58	68.98±0.93	50.75±1.11	50.13±0.45	58.00±0.42
20news	57.95±0.57	59.23±0.72	59.52±0.72	60.25±0.48	55.38±0.17	58.50±0.11	64.14±1.95	64.01±0.79	61.86±0.96	50.00±0.28	50.00±0.28	60.06±0.42
20news	57.43±0.39	57.59±0.72	57.59±0.72	57.89±0.09	56.85±0.16	57.32±0.11	68.75±0.91	65.60±1.74	65.60±1.74	50.00±0.28	50.00±0.28	60.06±0.42

Table 14.1: 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
aloi	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	<u>6.02$\pm$0.00</u>	6.07$\pm$0.00	6.09 $\pm$ 0.00	6.13 $\pm$ 0.00	6.15$\pm$0.00	6.09 $\pm$ 0.00
amazon	11.68 $\pm$ 0.00	11.68 $\pm$ 0.00	11.62 $\pm$ 0.00	11.61 $\pm$ 0.00	11.62 $\pm$ 0.00	11.65 $\pm$ 0.00	11.65 $\pm$ 0.00	11.69 $\pm$ 0.00	11.69 $\pm$ 0.00	11.60 $\pm$ 0.00	11.59 $\pm$ 0.00	11.65 $\pm$ 0.00
amthroid	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	66.11 $\pm$ 0.00	68.07$\pm$0.00	67.90 $\pm$ 0.00	67.26 $\pm$ 0.00	66.24 $\pm$ 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	29.58 $\pm$ 0.37	40.48 $\pm$ 0.81	46.70$\pm$1.22	37.36 $\pm$ 1.35	29.58 $\pm$ 0.58	24.37 $\pm$ 0.41	22.34 $\pm$ 0.53	32.06 $\pm$ 0.76
breasw	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	50.45$\pm$0.00	49.43 $\pm$ 0.00	49.64 $\pm$ 0.00	49.30 $\pm$ 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	79.35 $\pm$ 0.00	77.22 $\pm$ 0.00	78.33 $\pm$ 0.00	79.14 $\pm$ 0.00	80.67 $\pm$ 0.00	81.15$\pm$0.00	79.30 $\pm$ 0.00
cardiography	56.53 $\pm$ 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 $\pm$ 0.50	13.26 $\pm$ 0.61	14.50 $\pm$ 0.58	15.70 $\pm$ 0.65	16.10$\pm$0.68	14.31 $\pm$ 0.60
cinema	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 $\pm$ 0.42	21.36$\pm$0.76	21.22 $\pm$ 0.39	20.59 $\pm$ 0.33	20.00 $\pm$ 0.42	19.94 $\pm$ 0.44	20.62 $\pm$ 0.40
cover	63.67 $\pm$ 3.21	57.85 $\pm$ 3.52	51.55 $\pm$ 3.10	44.58 $\pm$ 2.49	42.11 $\pm$ 2.29	51.05 $\pm$ 2.42	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	82.34 $\pm$ 1.32	87.78 $\pm$ 0.99	89.44$\pm$0.96	88.33 $\pm$ 1.15	80.15 $\pm$ 1.33	73.68 $\pm$ 1.32	71.09 $\pm$ 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	61.10 $\pm$ 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 $\pm$ 5.58	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 $\pm$ 5.30	41.19 $\pm$ 5.73	37.35 $\pm$ 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 $\pm$ 2.85	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	78.95 $\pm$ 3.45	89.95 $\pm$ 1.80	91.10$\pm$4.41	69.28 $\pm$ 4.16	64.12 $\pm$ 5.59	64.25 $\pm$ 5.08	64.33 $\pm$ 6.16	70.62 $\pm$ 4.48
hup	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 $\pm$ 0.00	91.23 $\pm$ 1.42	91.44 $\pm$ 1.25	91.28 $\pm$ 1.36	94.39 $\pm$ 1.31
imdb	9.11 $\pm$ 0.00	9.09 $\pm$ 0.00	9.06 $\pm$ 0.00	9.07 $\pm$ 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
internets	52.20 $\pm$ 0.00	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$0.00	47.29 $\pm$ 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
landsat	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	46.49 $\pm$ 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	8.50 $\pm$ 0.00	8.50 $\pm$ 0.00	8.50 $\pm$ 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 $\pm$ 6.79	96.07 $\pm$ 6.79	95.68 $\pm$ 6.60	95.68 $\pm$ 6.60	96.16 $\pm$ 6.43	98.61 $\pm$ 1.02	98.43 $\pm$ 0.52	98.43 $\pm$ 0.52	98.70$\pm$0.83	98.57 $\pm$ 0.65	98.57 $\pm$ 0.65
magic gamma	86.30 $\pm$ 0.00	85.86 $\pm$ 0.00	85.28 $\pm$ 0.00	84.56 $\pm$ 0.00	84.29 $\pm$ 0.00	85.26 $\pm$ 0.00	85.86$\pm$0.00	83.25 $\pm$ 0.00	83.61 $\pm$ 0.00	83.61 $\pm$ 0.00	84.50 $\pm$ 0.00	84.50 $\pm$ 0.00
mammography	42.14 $\pm$ 0.00	41.51 $\pm$ 0.00	40.67 $\pm$ 0.00	40.37 $\pm$ 0.00	40.50 $\pm$ 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	39.83 $\pm$ 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	68.60 $\pm$ 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
optdigits	34.44 $\pm$ 0.00	30.53 $\pm$ 0.00	26.28 $\pm$ 0.00	22.67 $\pm$ 0.00	21.61 $\pm$ 0.00	27.11 $\pm$ 0.00	29.11$\pm$0.00	24.76 $\pm$ 0.00	21.10 $\pm$ 0.00	17.68 $\pm$ 0.00	16.62 $\pm$ 0.00	21.85 $\pm$ 0.00
pagesblocks	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	67.74 $\pm$ 0.00	67.87 $\pm$ 0.00	66.41 $\pm$ 0.00	66.13 $\pm$ 0.00	67.15 $\pm$ 0.00
pedigree	97.68 $\pm$ 0.00	97.31 $\pm$ 0.00	96.28 $\pm$ 0.00	90.01 $\pm$ 0.00	86.69 $\pm$ 0.00	93.59 $\pm$ 0.00	96.99$\pm$0.00	95.65 $\pm$ 0.00	81.40 $\pm$ 0.00	70.28 $\pm$ 0.00	67.39 $\pm$ 0.00	82.34 $\pm$ 0.00
planet	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	76.68 $\pm$ 2.22	75.66$\pm$2.91	73.62 $\pm$ 2.59	73.42 $\pm$ 2.98	73.71 $\pm$ 2.79	71.63 $\pm$ 2.86	74.01 $\pm$ 2.75
satellite	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	84.05 $\pm$ 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
satimage-2	96.10 $\pm$ 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	96.09$\pm$0.00	97.21 $\pm$ 0.00	97.39 $\pm$ 0.00	97.42$\pm$0.00	97.42 $\pm$ 0.00	97.22 $\pm$ 0.00
shuttle	99.16 $\pm$ 0.00	98.76 $\pm$ 0.00	98.22 $\pm$ 0.00	98.78 $\pm$ 0.00	98.77 $\pm$ 0.00	98.84 $\pm$ 0.00	97.86$\pm$0.00	97.34 $\pm$ 0.00	97.28 $\pm$ 0.00	97.22 $\pm$ 0.00	97.20 $\pm$ 0.00	97.38 $\pm$ 0.00
skin	56.70 $\pm$ 7.16	54.77 $\pm$ 7.80	54.75 $\pm$ 7.81	54.76 $\pm$ 7.81	48.78 $\pm$ 10.23	53.94 $\pm$ 7.83	98.31$\pm$0.34	50.26 $\pm$ 5.73	50.18 $\pm$ 5.75	50.33 $\pm$ 5.85	50.41$\pm$5.72	50.27 $\pm$ 5.76
snip	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
spambase	3.02 $\pm$ 0.00	2.89 $\pm$ 0.00	2.70 $\pm$ 0.00	2.76 $\pm$ 0.00	2.70 $\pm$ 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
speech	82.50 $\pm$ 3.71	77.11 $\pm$ 4.30	69.99 $\pm$ 5.29	65.85 $\pm$ 6.16	64.64 $\pm$ 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
stamps	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	80.94$\pm$0.00	81.09 $\pm$ 0.00	81.50 $\pm$ 0.00	81.90 $\pm$ 0.00	81.93$\pm$0.00	81.47 $\pm$ 0.00
thyroid	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
verbal	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
vowels	26.96 $\pm$ 0.00	26.71 $\pm$ 0.00	25.68 $\pm$ 0.00	24.67 $\pm$ 0.00	24.49 $\pm$ 0.00	25.70 $\pm$ 0.00	27.00$\pm$0.00	25.82 $\pm$ 0.00	23.77 $\pm$ 0.00	24.13 $\pm$ 0.00	23.87 $\pm$ 0.00	24.92 $\pm$ 0.00
waveform	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 $\pm$ 4.57	89.48$\pm$5.58	89.07 $\pm$ 3.41	88.67 $\pm$ 4.37	91.70 $\pm$ 3.59	91.82$\pm$3.52	90.15 $\pm$ 3.97
wbc	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	83.72 $\pm$ 5.56	82.05 $\pm$ 5.59	82.37 $\pm$ 4.91	82.33 $\pm$ 4.34	83.17 $\pm$ 5.03
wdbc	13.43 $\pm$ 0.00	12.36 $\pm$ 0.00	11.30 $\pm$ 0.00	10.40 $\pm$ 0.00	10.12 $\pm$ 0.00	11.52 $\pm$ 0.00	12.25					

Table 14.2: 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 across different HP settings.

dataset	ICL-0.1	ICL-0.01	ICL-0.001	ICL-0.0001	ICL-1e-05	ICL-avr	DTE-C-5	DTE-C-10	DTE-C-20	DTE-C-40	DTE-C-50	DTE-C-avr
aloi	5.50±0.09	5.39±0.07	5.50±0.08	5.59±0.04	5.46±0.01	5.49±0.02	5.76±0.03	5.82±0.02	5.72±0.05	5.73±0.05	5.91±0.00	5.79±0.01
amazon	10.06±0.10	10.08 ±0.03	9.91±0.13	10.01±0.05	9.99±0.02	10.01±0.05	11.01±0.43	10.99±1.11	11.05 ±0.93	11.05±0.93	9.52±0.00	10.42±0.00
amthroid	58.53 ±13.33	39.69±5.23	53.66±3.08	55.94±2.70	57.63±1.03	52.34±0.63	82.46±0.50	83.27±0.63	83.27 ±0.63	81.52±1.18	13.81±0.00	68.86±0.20
backdoor	85.81±2.45	88.14±1.38	88.99 ±1.12	88.96±1.20	86.24±0.70	87.63±1.03	43.90±3.90	61.21±2.58	63.64 ±1.61	4.83±0.09	4.83±0.09	35.68±0.79
breastw	98.65±0.81	98.46±0.51	98.98 ±0.57	98.50±0.42	95.32±1.11	97.98±0.30	88.69±0.50	89.49±0.20	94.04±1.49	98.56±0.68	99.22 ±0.11	94.00±0.69
campaign	47.46±0.41	44.03±2.05	47.94±0.25	49.22±0.91	51.41 ±0.51	48.01±0.47	49.90±2.01	46.77±1.41	48.40±1.60	20.25±0.00	37.11±0.64	37.11±0.64
cardio	48.31±14.30	65.70 ±11.26	63.55±5.07	61.75±2.55	29.67±1.62	53.90±3.32	69.32±0.43	70.25±0.47	17.78±0.59	17.55±0.00	49.02±0.29	49.02±0.29
cardiotocography	43.21±5.03	52.56 ±1.14	50.69±2.12	49.11±1.29	39.79±1.63	47.05±1.34	53.83±0.94	53.56±0.73	49.12±4.17	36.12±0.00	36.12±0.00	45.75±0.91
celeba	12.89±1.17	13.90 ±1.49	13.09±1.73	13.50±1.15	12.97±0.55	13.27±0.47	15.16±0.13	13.87±0.59	12.60±1.22	4.31±0.09	10.05±0.49	10.05±0.49
census	20.29±1.27	20.38±2.24	23.32±0.50	23.68 ±0.70	22.37±0.59	22.01±0.42	18.39±0.83	17.28±0.52	17.45±0.64	11.66±0.20	11.66±0.20	15.29±0.31
cover	22.58±13.63	10.32±4.93	33.04±2.12	47.50 ±19.76	40.81±22.18	31.42±3.27	71.28±4.54	66.66±3.73	38.26±3.29	19.4±0.07	19.4±0.07	35.03±0.82
donors	39.48±10.00	77.81±8.28	82.59±7.03	91.60±4.15	92.24 ±1.15	76.75±3.27	76.07±1.84	66.66±3.73	53.74±2.07	18.86±15.32	11.20±0.14	45.30±2.85
fault	65.37 ±3.40	64.58±1.81	63.83±0.74	62.41±0.30	72.24 ±0.96	64.33±1.09	63.92 ±1.40	62.97±0.40	63.62±1.06	51.69±0.25	51.74±0.32	58.79±0.41
fraud	51.45±12.34	49.97±0.27	56.24±9.47	62.41±10.30	72.24 ±0.96	58.46±1.58	68.85 ±0.35	67.17±4.68	24.97±21.19	0.34±0.03	0.34±0.03	30.33±4.25
glass	49.20±15.55	69.98±13.15	88.19 ±6.00	87.25±7.44	87.79±8.75	76.58±1.38	48.09 ±0.97	36.80±5.13	27.63±1.55	20.53±3.60	2.83±1.03	27.78±2.62
hepatitis	99.85±0.30	97.89±3.11	99.79±0.41	98.94±2.13	99.93 ±0.14	97.56±1.20	96.64 ±1.76	95.49±4.81	96.17±4.38	54.16±11.62	27.45±1.71	74.38±1.79
hup	94.83±9.03	95.76±3.25	97.86±1.72	99.77±0.25	99.53±0.60	97.56±1.20	96.64 ±1.76	95.49±4.81	96.17±4.38	54.16±11.62	27.45±1.71	74.38±1.79
indb	10.09±0.09	10.02±0.07	10.16±0.04	10.18±0.05	10.41 ±0.04	10.17±0.03	60.45 ±4.51	56.24±2.74	57.78±6.93	31.53±0.00	31.53±0.00	47.51±1.92
internucl	57.32±2.19	60.94±1.44	62.86±1.88	63.83 ±0.88	47.64±2.20	58.52±1.04	60.45 ±4.51	56.24±2.74	57.78±6.93	31.53±0.00	31.53±0.00	47.51±1.92
ionosphere	96.94±2.27	97.20±2.57	99.00 ±0.36	98.91±0.48	97.54±1.39	97.92±1.11	96.52±0.77	96.49±1.77	96.08 ±0.68	86.50±4.57	36.47±15.94	86.53±4.28
landsat	58.66 ±1.95	56.49±1.07	55.72±1.20	55.69±0.77	52.73±0.98	55.86±0.60	36.22 ±0.76	36.06±2.33	34.01±2.31	34.27±0.09	34.32±0.00	34.98±0.81
letter	11.22±1.58	10.84±0.96	9.41±0.34	9.71±0.55	15.87 ±1.71	11.41±0.43	8.92±0.10	9.30±0.23	8.96±0.09	11.76 ±0.00	11.76±0.00	10.09±0.04
lymphography	100.00±0.00	100.00±0.00	100.00±0.00	100.00±0.00	100.00±0.00	100.00±0.00	87.03±7.20	93.30±2.48	83.45±12.01	72.50±16.37	7.80±0.50	68.82±5.96
magic gamma	73.92±3.66	81.23±2.83	82.64±0.71	82.47±0.98	83.45 ±1.83	80.74±1.03	89.01±0.45	89.46 ±0.29	89.15±4.44	66.79±18.08	52.03±0.00	77.29±3.52
mnist	33.51±8.47	37.72 ±7.54	27.06±2.69	26.94±1.92	17.30±1.14	28.57±1.33	35.02±2.57	32.22±4.57	35.30±2.06	4.54±0.00	28.91±1.92	28.91±1.92
mnistmn	49.27±1.39	50.40±3.02	54.46±1.34	63.49±1.45	65.20 ±1.41	59.48±1.94	60.64 ±1.13	56.20±2.31	55.17±3.05	16.86±0.00	16.86±0.00	41.15±1.11
mnistmn	49.27±1.39	50.40±3.02	54.46±1.34	63.49±1.45	65.20 ±1.41	59.48±1.94	60.64 ±1.13	56.20±2.31	55.17±3.05	16.86±0.00	16.86±0.00	41.15±1.11
mnistmn	49.27±1.39	50.40±3.02	54.46±1.34	63.49±1.45	65.20 ±1.41	59.48±1.94	60.64 ±1.13	56.20±2.31	55.17±3.05	16.86±0.00	16.86±0.00	41.15±1.11
mnistmn	49.27±1.39	50.40±3.02	54.46±1.34	63.49±1.45	65.20 ±1.41	59.48±1.94	60.64 ±1.13	56.20±2.31	55.17±3.05	16.86±0.00	16.86±0.00	41.15±1.11
mnistmn	49.27±1.39	50.40±3.02	54.46±1.34	63.49±1.45	65.20 ±1.41	59.48±1.94	60.64 ±1.13	56.20±2.31	55.17±3.05	16.86±0.00	16.86±0.00	41.15±1.11
mnistmn	49.27±1.39	50.40±3.02	54.46±1.34	63.49±1.45	65.20 ±1.41	59.48±1.94	60.64 ±1.13	56.20±2.31	55.17±3.05	16.86±0.00	16.86±0.00	41.15±1.11
mnistmn	49.27±1.39	50.40±3.02	54.46±1.34	63.49±1.45	65.20 ±1.41	59.48±1.94	60.64 ±1.13	56.20±2.31	55.17±3.05	16.86±0.00	16.86±0.00	41.15±1.11
mnistmn	49.27±1.39	50.40±3.02	54.46±1.34	63.49±1.45	65.20 ±1.41	59.48±1.94	60.64 ±1.13	56.20±2.31	55.17±3.05	16.86±0.00	16.86±0.00	41.15±1.11
mnistmn	49.27±1.39	50.40±3.02	54.46±1.34	63.49±1.45	65.20 ±1.41	59.48±1.94	60.64 ±1.13	56.20±2.31	55.17±3.05	16.86±0.00	16.86±0.00	41.15±1.11
mnistmn	49.27±1.39	50.40±3.02	54.46±1.34	63.49±1.45	65.20 ±1.41	59.48±1.94	60.64 ±1.13	56.20±2.31	55.17±3.05	16.86±0.00	16.86±0.00	41.15±1.11
mnistmn	49.27±1.39	50.40±3.02	54.46±1.34	63.49±1.45	65.20 ±1.41	59.48±1.94	60.64 ±1.13	56.20±2.31	55.17±3.05	16.86±0.00	16.86±0.00	41.15±1.11
mnistmn	49.27±1.39	50.40±3.02	54.46±1.34	63.49±1.45	65.20 ±1.41	59.48±1.94	60.64 ±1.13	56.20±2.31	55.17±3.05	16.86±0.00	16.86±0.00	41.15±1.11
mnistmn	49.27±1.39	50.40±3.02	54.46±1.34	63.49±1.45	65.20 ±1.41	59.48±1.94	60.64 ±1.13	56.20±2.31	55.17±3.05	16.86±0.00	16.86±0.00	41.15±1.11
mnistmn	49.27±1.39	50.40±3.02	54.46±1.34	63.49±1.45	65.20 ±1.41	59.48±1.94	60.64 ±1.13	56.20±2.31	55.17±3.05	16.86±0.00	16.86±0.00	41.15±1.11
mnistmn	49.27±1.39	50.40±3.02	54.46±1.34	63.49±1.45	65.20 ±1.41	59.48±1.94	60.64 ±1.13	56.20±2.31	55.17±3.05	16.86±0.00	16.86±0.00	41.15±1.11
mnistmn	49.27±1.39	50.40±3.02	54.46±1.34	63.49±1.45	65.20 ±1.41	59.48±1.94	60.64 ±1.13	56.20±2.31	55.17±3.05	16.86±0.00	16.86±0.00	41.15±1.11
mnistmn	49.27±1.39	50.40±3.02	54.46±1.34	63.49±1.45	65.20 ±1.41	59.48±1.94	60.64 ±1.13	56.20±2.31	55.17±3.05	16.86±0.00	16.86±0.00	41.15±1.11
mnistmn	49.27±1.39	50.40±3.02	54.46±1.34	63.49±1.45	65.20 ±1.41	59.48±1.94	60.64 ±1.13	56.20±2.31	55.17±3.05	16.86±0.00	16.86±0.00	41.15±1.11
mnistmn	49.27±1.39	50.40±3.02	54.46±1.34	63.49±1.45	65.20 ±1.41	59.48±1.94	60.64 ±1.13	56.20±2.31	55.17±3.05	16.86±0.00	16.86±0.00	41.15±1.11
mnistmn	49.27±1.39	50.40±3.02	54.46±1.34	63.49±1.45	65.20 ±1.41	59.48±1.94	60.64 ±1.13	56.20±2.31	55.17±3.05	16.86±0.00	16.86±0.00	41.15±1.11
mnistmn	49.27±1.39	50.40±3.02	54.46±1.34	63.49±1.45	65.20 ±1.41	59.48±1.94	60.64 ±1.13	56.20±2.31	55.17±3.05	16.86±0.00	16.86±0.00	41.15±1.11
mnistmn	49.27±1.39	50.40±3.02	54.46±1.34	63.49±1.45	65.20 ±1.41	59.48±1.94	60.64 ±1.13	56.20±2.31	55.17±3.05	16.86±0.00	16.86±0.00	41.15±1.11
mnistmn	49.27±1.39	50.40±3.02	54.46±1.34	63.49±1.45	65.20 ±1.41	59.48±1.94	60.64 ±1.13	56.20±2.31	55.17±3.05	16.86±0.00	16.86±0.00	41.15±1.11
mnistmn	49.27±1.39	50.40±3.02	54.46±1.34	63.49±1.45	65.20 ±1.41	59.48±1.94	60.64 ±1.13	56.20±2.31	55.17±3.05	16.86±0.00	16.86±0.00	41.15±1.11
mnistmn	49.27±1.39	50.40±3.02	54.46±1.34	63.49±1.45	65.20 ±1.41	59.48±1.94	60.64 ±1.13	56.20±2.31	55.17±3.05	16.86±0.00	16.86±0.00	41.15±1.11
mnistmn	49.27±1.39	50.40±3.02	54.46±1.34	63.49±1.45	65.20 ±1.41	59.48±1.94	60.64 ±1.13	56.20±2.31	55.17±3.05	16.86±0.00	16.86±0.00	41.15±1.11
mnistmn	49.27±1.39	50.40±3.02	54.46±1.34	63.49±1.45	65.20 ±1.41	59.48±1.94	60.64 ±1.13	56.20±2.31	55.17±3.05	16.86±0.00	16.86±0.00	41.15±1.11
mnistmn	49.27±1.39	50.40±3.02	54.46±1.34	63.49±1.45	65.20 ±1.41	59.48±1.94	60.64 ±1.13	56.20±2.31	55.17±3.05	16.86±0.00	16.86±0.00	41.15±1.11
mnistmn	49.27±1.39	50.40±3.02	54.46±1.34	63.49±1.45	65.20 ±1.41	59.48±1.94	60.64 ±1.13	56.20±2.31	55.17±3.05	16.86±0.00	16.86±0.00	41.15±1.11
mnistmn	49.27±1.39	50.40±3.02	54.46±1.34	63.49±1.45	65.20 ±1.41	59.48±1.94	60.64 ±1.13	56.20±2.31	55.17±3.05	16.86±0.00	16.86±0.00	41.15±1.11
mnistmn	49.27±1.39	50.40±3.02	54.									

Table 15.1: 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	5.70 $\pm$ 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	5.67 $\pm$ 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	11.00 $\pm$ 0.00	10.60 $\pm$ 0.00	11.40 $\pm$ 0.00	10.20 $\pm$ 0.00	10.60 $\pm$ 0.00	11.20 $\pm$ 0.00	11.20 $\pm$ 0.00	10.92 $\pm$ 0.00
amthyroid	62.55 $\pm$ 0.00	61.80 $\pm$ 0.00	60.67 $\pm$ 0.00	58.99 $\pm$ 0.00	58.99 $\pm$ 0.00	60.86 $\pm$ 0.00	61.99 $\pm$ 0.00	60.49 $\pm$ 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 $\pm$ 1.04	52.30 $\pm$ 1.87	40.62 $\pm$ 1.46	30.25 $\pm$ 1.34	26.96 $\pm$ 1.20	42.56 $\pm$ 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	95.99 $\pm$ 0.39	95.99 $\pm$ 0.39	96.22 $\pm$ 0.43	96.00 $\pm$ 0.44	96.05 $\pm$ 0.43	95.99 $\pm$ 0.39	95.99 $\pm$ 0.39	95.99 $\pm$ 0.32	95.97 $\pm$ 0.31
campaign	49.94 $\pm$ 0.00	50.62 $\pm$ 0.00	51.14 $\pm$ 0.00	51.38 $\pm$ 0.00	51.57 $\pm$ 0.00	50.93 $\pm$ 0.00	50.37 $\pm$ 0.00	50.86 $\pm$ 0.00	51.27 $\pm$ 0.00	51.70 $\pm$ 0.00	51.29 $\pm$ 0.00	51.10 $\pm$ 0.00
cardio	63.64 $\pm$ 0.00	61.36 $\pm$ 0.00	61.93 $\pm$ 0.00	64.36 $\pm$ 0.00	64.20 $\pm$ 0.00	62.95 $\pm$ 0.00	61.93 $\pm$ 0.00	61.93 $\pm$ 0.00	64.20 $\pm$ 0.00	67.61 $\pm$ 0.00	69.32 $\pm$ 0.00	65.00 $\pm$ 0.00
cardiotocography	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	46.35 $\pm$ 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	22.23 $\pm$ 0.42	18.41 $\pm$ 0.65	19.30 $\pm$ 0.81	20.27 $\pm$ 0.68	20.48 $\pm$ 0.68	19.11 $\pm$ 0.61
cinema	22.22 $\pm$ 0.54	21.93 $\pm$ 0.52	21.46 $\pm$ 0.25	21.38 $\pm$ 0.48	21.12 $\pm$ 0.29	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	39.92 $\pm$ 1.99	52.19 $\pm$ 1.92
donors	97.17 $\pm$ 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	76.62 $\pm$ 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	57.26 $\pm$ 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 $\pm$ 3.53	49.29 $\pm$ 3.24	46.54 $\pm$ 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 $\pm$ 4.21	18.37 $\pm$ 2.40	17.81 $\pm$ 2.98	29.42 $\pm$ 2.61	29.87 $\pm$ 9.62	22.53 $\pm$ 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	75.95 $\pm$ 4.78	71.99 $\pm$ 4.14	84.54 $\pm$ 4.50	81.98 $\pm$ 4.50	66.04 $\pm$ 4.52	62.31 $\pm$ 6.09	60.39 $\pm$ 6.21	60.04 $\pm$ 7.30	66.15 $\pm$ 4.71
hup	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	98.57 $\pm$ 2.86	92.67 $\pm$ 0.91	92.67 $\pm$ 0.91	92.67 $\pm$ 0.91	95.32 $\pm$ 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	5.28 $\pm$ 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
internets	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	49.57 $\pm$ 0.00	51.90 $\pm$ 0.00	46.20 $\pm$ 0.00	45.11 $\pm$ 0.00	45.11 $\pm$ 0.00	45.11 $\pm$ 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	91.12 $\pm$ 1.24	90.23 $\pm$ 1.86	89.23 $\pm$ 1.86	89.45 $\pm$ 1.91	85.06 $\pm$ 3.37	82.73 $\pm$ 3.12	87.86 $\pm$ 1.86
landsat	52.29 $\pm$ 0.00	51.24 $\pm$ 0.00	49.06 $\pm$ 0.00	45.99 $\pm$ 0.00	45.39 $\pm$ 0.00	48.79 $\pm$ 0.00	51.46 $\pm$ 0.00	49.29 $\pm$ 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 $\pm$ 4.21	93.61 $\pm$ 7.97	94.67 $\pm$ 6.53	92.71 $\pm$ 6.09	91.66 $\pm$ 7.34	94.11 $\pm$ 5.93	91.57 $\pm$ 5.96	90.69 $\pm$ 3.65	90.69 $\pm$ 3.65	92.96 $\pm$ 4.97	92.96 $\pm$ 4.97	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	74.75 $\pm$ 0.00	75.61 $\pm$ 0.00	76.17 $\pm$ 0.00	73.13 $\pm$ 0.00	74.60 $\pm$ 0.00	73.49 $\pm$ 0.00	73.00 $\pm$ 0.00	74.48 $\pm$ 0.00
mammography	41.92 $\pm$ 0.00	41.15 $\pm$ 0.00	44.23 $\pm$ 0.00	44.23 $\pm$ 0.00	44.23 $\pm$ 0.00	43.15 $\pm$ 0.00	40.38 $\pm$ 0.00	43.46 $\pm$ 0.00	43.85 $\pm$ 0.00	43.46 $\pm$ 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	69.86 $\pm$ 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
optdigits	30.00 $\pm$ 0.00	24.00 $\pm$ 0.00	12.00 $\pm$ 0.00	7.33 $\pm$ 0.00	6.67 $\pm$ 0.00	16.00 $\pm$ 0.00	21.33 $\pm$ 0.00	12.00 $\pm$ 0.00	6.67 $\pm$ 0.00	2.67 $\pm$ 0.00	2.00 $\pm$ 0.00	8.93 $\pm$ 0.00
pagelabels	59.41 $\pm$ 0.00	59.22 $\pm$ 0.00	59.61 $\pm$ 0.00	58.43 $\pm$ 0.00	58.43 $\pm$ 0.00	59.02 $\pm$ 0.00	59.02 $\pm$ 0.00	59.22 $\pm$ 0.00	56.08 $\pm$ 0.00	56.27 $\pm$ 0.00	56.27 $\pm$ 0.00	58.16 $\pm$ 0.00
pendigits	94.23 $\pm$ 0.00	92.31 $\pm$ 0.00	91.03 $\pm$ 0.00	80.13 $\pm$ 0.00	78.21 $\pm$ 0.00	87.18 $\pm$ 0.00	90.38 $\pm$ 0.00	90.38 $\pm$ 0.00	73.72 $\pm$ 0.00	62.82 $\pm$ 0.00	62.82 $\pm$ 0.00	76.41 $\pm$ 0.00
plot	74.73 $\pm$ 2.13	72.94 $\pm$ 2.46	71.48 $\pm$ 1.96	71.44 $\pm$ 2.36	71.93 $\pm$ 1.99	72.50 $\pm$ 1.97	71.03 $\pm$ 2.53	70.18 $\pm$ 2.12	71.58 $\pm$ 2.16	71.67 $\pm$ 2.10	72.03 $\pm$ 2.02	71.30 $\pm$ 2.00
satellite	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	70.71 $\pm$ 0.00	71.81 $\pm$ 0.00	70.73 $\pm$ 0.00	69.84 $\pm$ 0.00	67.93 $\pm$ 0.00	67.29 $\pm$ 0.00	69.52 $\pm$ 0.00
satimage-2	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	90.14 $\pm$ 0.00	91.35 $\pm$ 0.00	92.96 $\pm$ 0.00	92.96 $\pm$ 0.00	92.96 $\pm$ 0.00	92.11 $\pm$ 0.00
shuttle	96.35 $\pm$ 0.00	96.23 $\pm$ 0.00	96.15 $\pm$ 0.00	98.12 $\pm$ 0.00	98.12 $\pm$ 0.00	98.19 $\pm$ 0.00	96.23 $\pm$ 0.00	98.06 $\pm$ 0.00	98.06 $\pm$ 0.00	98.09 $\pm$ 0.00	98.09 $\pm$ 0.00	98.11 $\pm$ 0.00
skin	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	95.83 $\pm$ 0.17	94.56 $\pm$ 0.29	91.73 $\pm$ 0.16	90.60 $\pm$ 0.28	93.89 $\pm$ 0.13
snip	68.05 $\pm$ 5.12	68.05 $\pm$ 5.12	69.59 $\pm$ 3.95	69.59 $\pm$ 3.95	63.34 $\pm$ 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
spambase	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	78.98 $\pm$ 0.00	78.98 $\pm$ 0.00	78.80 $\pm$ 0.00	79.48 $\pm$ 0.00
speech	3.28 $\pm$ 0.00	3.28 $\pm$ 0.00	1.64 $\pm$ 0.00	1.64 $\pm$ 0.00	3.28 $\pm$ 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
stamps	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	78.57 $\pm$ 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
thyroid	75.27 $\pm$ 0.00	75.27 $\pm$ 0.00	74.19 $\pm$ 0.00	74.19 $\pm$ 0.00	74.62 $\pm$ 0.00	74.62 $\pm$ 0.00	75.27 $\pm$ 0.00	73.12 $\pm$ 0.00	73.12 $\pm$ 0.00	73.12 $\pm$ 0.00	73.12 $\pm$ 0.00	73.55 $\pm$ 0.00
verbal	49.03 $\pm$ 4.93	32.73 $\pm$ 5.11	24.06 $\pm$ 3.90	15.07 $\pm$ 1.64	12.51 $\pm$ 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
vowels	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	28.80 $\pm$ 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	26.80 $\pm$ 0.00
waveform	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
wbc	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
wdbc	85.05 $\pm$ 6.11</											

Table 15.2: 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-avr	DTE-C-5	DTE-C-10	DTE-C-20	DTE-C-40	DTE-C-50	DTE-C-avg
aloi	4.51±0.69	4.34±0.42	5.28±0.47	4.68±0.30	4.16±0.38	4.94±0.07	4.75±0.27	4.27±0.19	4.28±0.10	4.51±0.17	0.00±0.00	3.56±0.03
amazon	10.44±0.46	9.76±0.84	9.92±0.84	10.08±0.35	9.52±0.43	9.94±0.32	11.48±0.97	11.96±1.68	11.60±2.07	0.00±0.00	0.00±0.00	7.50±0.72
amthroid	54.87±13.24	42.25±3.55	53.45±5.45	57.53±2.91	52.56±3.29	54.72±5.45	77.94±0.28	75.43±0.93	77.53±0.25	75.43±0.93	0.00±0.00	61.63±0.20
backdoor	87.17±0.98	87.32±0.99	87.11±1.09	86.85±0.95	85.37±1.01	86.76±1.00	46.19±0.25	83.03±2.14	84.50±0.60	0.00±0.00	0.00±0.00	42.75±1.54
breasw	95.98±0.34	96.07±0.94	96.80±0.40	97.44±0.55	96.48±0.28	96.48±0.28	88.50±1.59	90.10±1.35	92.46±1.78	95.31±0.70	96.11±0.44	92.50±0.79
campaign	48.12±0.36	46.81±1.72	50.68±0.66	51.37±0.85	53.40±0.51	50.07±0.49	51.98±0.70	52.45±1.07	52.33±1.00	0.00±0.00	0.00±0.00	31.35±0.44
cardio	49.09±11.28	61.93±5.57	58.86±1.59	57.95±2.30	40.57±4.28	53.68±1.27	58.30±0.58	57.84±0.43	57.07±0.23	0.34±0.68	0.00±0.00	34.91±0.09
cardiography	36.14±1.28	41.07±1.73	39.18±4.49	35.36±2.14	32.66±1.59	36.88±1.27	39.91±1.05	39.48±1.54	37.23±1.73	62.02±0.00	62.02±0.00	48.23±0.39
celeba	15.42±2.29	17.97±2.55	17.20±1.92	17.46±1.17	16.17±0.65	16.84±0.88	19.18±2.74	17.12±1.45	14.31±2.28	0.00±0.00	0.00±0.00	10.12±1.04
census	22.72±1.73	24.06±2.05	25.80±1.34	27.15±1.12	24.06±1.31	24.76±0.50	17.58±1.42	17.54±1.79	16.44±1.41	0.00±0.00	0.00±0.00	10.31±0.62
cover	26.77±15.24	15.53±8.17	42.68±17.00	53.70±16.88	44.34±20.24	36.61±2.61	76.51±2.37	68.92±4.22	46.54±3.93	0.00±0.00	0.00±0.00	38.39±0.62
donors	43.71±11.52	81.85±3.56	83.59±4.47	89.28±2.66	92.77±1.39	78.24±2.61	77.09±1.87	75.05±1.18	63.08±1.70	11.80±23.60	96.91±0.24	47.58±4.61
fault	60.33±11.51	59.52±2.09	58.57±1.11	58.37±0.79	58.87±1.51	59.13±1.36	55.57±1.57	53.05±0.81	55.33±1.66	97.03±0.60	96.91±0.24	72.00±0.47
glass	57.34±10.13	48.97±8.02	58.90±6.77	67.24±5.84	82.50±5.41	70.87±5.12	75.61±7.76	54.25±5.90	22.55±24.73	0.00±0.00	0.00±0.00	30.48±4.66
frass	43.53±20.21	57.05±16.03	84.05±6.11	87.24±5.04	79.18±5.41	68.65±2.11	75.61±7.76	54.25±5.90	22.55±24.73	0.00±0.00	0.00±0.00	30.48±4.66
hepatitis	99.64±0.71	94.69±0.71	99.64±0.71	99.64±0.71	99.64±0.71	98.65±2.11	96.40±3.01	94.63±4.90	51.68±9.78	18.97±10.60	0.00±0.00	70.84±3.40
hip	93.91±10.35	96.07±3.07	97.69±1.51	99.36±0.19	99.14±0.43	97.23±2.45	38.08±11.97	50.80±3.02	18.33±10.68	16.75±12.06	0.00±0.00	24.79±12.88
imdb	10.52±0.65	10.44±0.54	9.84±0.37	9.76±0.15	10.40±0.33	10.19±0.23	6.64±0.89	8.40±1.23	7.32±1.04	0.00±0.00	0.00±0.00	4.47±0.43
internats	55.92±2.66	57.45±0.61	57.77±1.10	58.26±0.50	49.02±1.45	55.68±0.94	67.99±1.98	63.87±0.95	64.95±2.24	41.85±0.00	0.00±0.00	56.50±0.87
ionosphere	92.64±4.66	91.41±4.67	93.86±1.63	94.48±0.71	94.49±1.56	93.38±2.21	89.67±1.44	89.41±1.53	89.52±1.13	78.12±2.07	49.44±16.82	79.23±4.14
landsat	49.50±1.24	47.94±1.88	54.51±0.52	54.25±0.74	47.97±1.09	50.83±0.68	35.45±0.79	35.47±3.76	31.67±3.98	50.29±8.34	54.46±0.00	41.47±2.39
letter	6.80±4.21	4.00±1.55	3.60±1.36	3.20±0.75	11.60±2.06	5.84±1.11	2.41±1.30	3.00±1.55	3.20±1.17	0.00±0.00	0.00±0.00	1.72±0.37
lymphography	100.00±0.00	100.00±0.00	100.00±0.00	100.00±0.00	100.00±0.00	100.00±0.00	77.11±6.79	79.84±1.53	74.75±5.79	60.02±26.09	0.00±0.00	58.34±6.75
magic gamma	64.88±2.50	69.99±2.98	69.99±0.87	70.21±0.50	71.64±0.93	69.34±1.08	79.82±0.48	80.94±0.33	80.19±0.73	89.73±7.81	96.10±0.00	85.36±1.69
mammography	36.62±8.76	38.15±8.41	27.69±1.80	29.08±2.63	18.15±1.23	29.94±2.52	32.69±5.73	34.62±2.19	35.08±2.68	39.69±2.19	0.00±0.00	28.42±0.00
mnist	45.37±1.84	46.20±3.18	50.54±1.26	59.20±1.30	62.66±2.09	52.79±0.77	63.86±1.77	56.74±2.71	53.74±3.75	0.00±0.00	0.00±0.00	34.87±1.48
mnist	100.00±0.00	100.00±0.00	100.00±0.00	100.00±0.00	100.00±0.00	100.00±0.00	100.00±0.00	100.00±0.00	100.00±0.00	0.00±0.00	0.00±0.00	60.00±0.00
mnist	29.87±5.68	41.20±3.49	44.67±5.56	71.73±0.53	62.00±5.11	46.69±2.00	14.80±1.81	6.40±3.88	9.33±5.72	0.00±0.00	0.00±0.00	6.11±1.52
optdigits	62.16±2.84	64.08±2.96	62.31±2.65	63.88±1.08	62.00±1.05	62.93±0.29	62.71±1.24	61.25±0.32	60.47±0.69	0.00±0.00	0.00±0.00	49.33±0.57
pageslocks	46.03±17.81	60.51±10.58	59.23±5.15	66.03±5.18	51.03±5.57	56.56±5.52	63.97±6.36	54.36±7.20	43.46±5.50	0.00±0.00	0.00±0.00	32.36±1.43
pendigits	68.77±4.96	68.29±2.03	74.95±1.27	71.40±1.67	68.83±0.58	70.45±1.78	66.14±2.35	63.89±4.28	64.71±3.39	67.99±2.34	65.96±4.44	65.74±2.10
pima	68.94±10.44	72.05±2.96	76.27±0.81	78.70±0.90	74.22±0.57	73.62±2.69	72.06±0.60	72.07±0.77	72.63±0.32	91.51±9.02	96.02±0.00	81.04±1.96
satellite	38.03±8.59	72.68±32.26	91.83±0.56	89.58±1.13	56.34±12.38	69.69±7.51	78.51±3.03	60.85±5.87	46.20±5.07	26.20±32.11	0.00±0.00	42.31±7.14
satimage-2	97.17±1.11	98.27±0.10	98.83±0.14	98.91±0.12	96.38±0.21	98.31±0.24	98.00±0.01	97.98±0.00	97.73±0.10	92.83±2.84	0.00±0.00	77.31±0.58
skin	38.03±28.70	72.09±8.03	67.99±9.92	54.29±12.30	49.74±13.81	56.43±9.55	82.15±0.37	82.23±0.37	81.66±0.57	80.48±0.35	54.74±0.80	76.25±0.32
snip	39.28±18.81	59.49±7.84	54.93±13.50	48.81±0.96	50.03±10.81	56.43±9.55	69.59±3.95	69.59±3.95	27.32±14.61	43.13±3.67	43.13±3.67	76.25±0.32
speech	76.97±2.12	76.88±3.76	80.35±0.48	80.23±0.62	74.93±0.44	77.87±0.46	80.19±0.18	80.27±0.41	80.56±0.29	87.61±0.00	87.61±0.00	83.25±0.07
stamps	2.95±1.23	3.61±1.23	2.95±2.41	4.59±1.61	3.48±0.92	3.48±0.92	3.93±0.80	2.95±1.61	3.28±1.80	0.00±0.00	0.00±0.00	2.03±0.25
thyroid	34.84±12.18	42.45±13.97	83.03±3.34	76.24±6.92	73.47±7.50	62.00±6.23	63.62±10.79	58.37±9.83	57.63±9.24	56.93±11.83	14.60±29.20	50.23±10.89
verbal	68.39±4.17	61.29±1.80	61.08±4.63	63.87±3.57	33.76±5.95	57.68±1.74	73.76±2.21	76.13±1.72	75.27±0.00	78.28±0.80	0.00±0.00	60.69±0.57
vowels	23.17±6.96	29.15±9.05	52.53±11.60	64.81±3.19	54.89±6.76	44.91±2.82	43.88±5.22	36.26±11.73	32.59±10.94	24.68±8.71	3.48±6.96	28.18±7.79
waveform	22.40±17.59	18.80±6.52	30.80±7.44	30.80±12.43	24.00±7.04	25.36±2.59	40.00±5.51	40.80±3.25	45.20±0.98	0.00±0.00	0.00±0.00	25.20±1.36
waveform	47.80±3.06	38.20±17.57	35.40±7.47	11.60±1.62	28.40±4.22	32.28±5.79	11.20±1.72	11.80±2.14	12.20±1.60	0.00±0.00	0.00±0.00	7.04±0.23
wbc	78.86±7.89	84.03±7.23	96.89±4.35	95.71±4.90	90.32±5.34	89.16±3.80	40.59±7.11	34.81±6.65	40.96±12.93	67.85±5.79	0.00±0.00	36.84±4.92
wdbc	64.32±18.21	80.42±5.68	84.10±6.72	87.13±4.26	78.33±5.25	78.86±2.34	78.62±5.69	70.31±4.12	77.59±7.50	10.43±20.87	0.00±0.00	47.39±3.91
wine	61.55±7.00	11.36±4.93	37.04±8.60	39.61±4.50	37.59±2.50	26.35±2.35	19.53±2.50	19.92±4.34	5.37±2.97	16.96±5.74	0.00±0.00	12.36±2.38
wine	97.95±4.09	96.72±3.06	96.25±2.23	98.18±1.64	97.67±1.33	97.76±2.03	98.86±1.44	95.85±5.74	97.27±6.94	14.96±5.74	0.00±0.00	66.58±4.34
wpc	81.91±5.11	63.78±10.33	88.98±0.38	88.64±0.69	83.46±1.33	81.35±3.33	49.23±1.27	59.07±4.32	56.98±1.54	36.48±1.91	38.96±14.97	58.93±5.76
wpbc	46.31±2.91	48.95±17.18	49.35±0.54	49.31±0.62	48.41±0.93	48.41±0.87	48.41±0.87	48.41±0.87	48.41±0.87	48.41±0.87	98.22±0.01	33.94±0.44
yelp	7.64±0.70	8.20±0.38	7.52±0.65	7.56±0.29	7.56±0.37	7.70±0.22	7.70±0.22	7.70±0.22	7.70±0.22	7.70±0.22	0.00±0.00	9.11±1.45
mnist-C	45.87±0.49	50.05±0.46	51.49±0.27	51.46±0.39	47.26±0.49	47.26±0.49	47.26±0.49	47.26±0.49	47.26±0.49	47.26±0.49	0.00±0.00	29.35±0.11
FashionMNIST	56.83±0.32	62.33±0.40	64.84±0.26	64.81±0.16	58.13±0.29	58.13±0.29	58.13±0.29	58.13±0.29	58.13±0.29	58.13±0.29	0.00±0.00	35.54±0.10
CIFAR10	16.33±0.58	20.51±0.66	22.84±0.26	22.87±0.57	15.56±0.45	15.56±0.45	15.56±0.45	15.56±0.45	15.56±0.45	15.56±0.45	0.00±0.00	14.40±0.12
SVHN	16.44±0.99	19.22±0.23	20.00±0.14	20.07±0.36	17.85±0.22	17.85±0.22	17.85±0.22	17.85±0.22	17.85±0.22	17.85±0.22	0.00±0.00	11.78±0.08
MVTeC-AD	79.81±0.59	81.35±0.82	82.39±0.82	82.30±0.46	77.51±0.66	77.51±0.66	77.51±0.66	77.51±0.66	77.51±0.66	77.51±0.66	0.00±0.00	71.97±0.92
20news	12.82±0.91	14.12±1.10	15.05±1.36	16.54±1.42	12.79±0.95	12.79±0.95	12.79±0.95	12.79±0.95	12.79±0.95	12.79±0.95	0.18±0.35	10.44±0.45
20news	14.17±0.24	14.47±0.24	14.47±0.12	14.53±0.20	13.69±0.12	14.26±0.73	19.22±2.15	16.99±0.93	19.22±2.15	19.22±2.15	0.00±0.00	12.53±0.63

Table 16.1: Average AUROC $\pm$ standard dev. over five seeds for the semi-supervised setting on ADBench. Rank of each model among 32 models (26 baselines + 4 variants of top-4 baselines + 2 FoMo-0D variants w/ $D = 100$ and $D = 20$) 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. Last four rows show avg_rank of methods across datasets, and p -values of the Wilcoxon signed rank test comparing FoMo-0D ($D = 100$) with other baselines. The previous four rows are the same for FoMo-0D ($D = 20$), when ranking 31 models (26 baselines + 4 variants of top-4 baselines + FoMo-0D w/ $D = 20$).

[illegible]

Table 16.2: Average AUROC $\pm$ 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.

Dataset	VAE	PCA	PlanerFlow	HDBS	GANomaly	GOAD	DIF	COPOD	ECOD	DeepSVDD	LODA	DAGMM	DROCC	DTE- $\mathcal{N}_{\text{opt}}$	KN $\mathcal{N}_{\text{opt}}$	ICL- $\mathcal{N}_{\text{opt}}$	DTE- $\mathcal{N}_{\text{opt}}$
aloi	54.04 $\pm$ 0.03(5)	54.04 $\pm$ 0.03(5)	48.52 $\pm$ 2.34(29)	52.23 $\pm$ 0.00(7)	53.94 $\pm$ 1.65(5)	48.01 $\pm$ 0.02(30)	51.14 $\pm$ 0.01(3)	49.31 $\pm$ 0.02(4)	51.73 $\pm$ 0.00(8)	50.89 $\pm$ 2.05(15)	49.24 $\pm$ 2.75(25)	50.84 $\pm$ 2.07(17)	50.01 $\pm$ 0.02(2)	51.25 $\pm$ 0.04(1)	51.61 $\pm$ 0.00(9)	47.63 $\pm$ 0.15(31)	50.29 $\pm$ 0.09(21)
amazon	85.44 $\pm$ 0.02(9)	85.19 $\pm$ 0.00(1)	93.19 $\pm$ 2.11(1)	66.02 $\pm$ 0.03(1)	87.17 $\pm$ 1.45(3)	81.01 $\pm$ 1.16(25)	87.36 $\pm$ 0.76(16)	76.72 $\pm$ 0.02(8)	78.45 $\pm$ 0.00(30)	91.01 $\pm$ 0.62(2)	77.35 $\pm$ 2.51(27)	72.33 $\pm$ 1.91(29)	88.93 $\pm$ 0.73(2)	92.64 $\pm$ 0.00(1)	91.31 $\pm$ 0.07(2)	84.22 $\pm$ 0.31(2)	88.08 $\pm$ 0.05(1)
backdoor	64.67 $\pm$ 0.04(20)	64.67 $\pm$ 0.04(20)	76.03 $\pm$ 1.11(56)(18)	70.81 $\pm$ 0.88(21)	87.17 $\pm$ 1.45(3)	81.01 $\pm$ 1.16(25)	87.36 $\pm$ 0.76(16)	76.72 $\pm$ 0.02(8)	78.45 $\pm$ 0.00(30)	91.01 $\pm$ 0.62(2)	77.35 $\pm$ 2.51(27)	72.33 $\pm$ 1.91(29)	88.93 $\pm$ 0.73(2)	92.64 $\pm$ 0.00(1)	91.31 $\pm$ 0.07(2)	84.22 $\pm$ 0.31(2)	88.08 $\pm$ 0.05(1)
breast	99.22 $\pm$ 0.00(1)	99.22 $\pm$ 0.00(1)	97.93 $\pm$ 0.82(21)	97.06 $\pm$ 0.17(6)	99.21 $\pm$ 0.73(5)	87.86 $\pm$ 0.33(15)	56.34 $\pm$ 1.19(31)	99.46 $\pm$ 0.09(3)	99.14 $\pm$ 0.22(11)	96.96 $\pm$ 0.92(23)	98.13 $\pm$ 0.35(21)	89.53 $\pm$ 0.12(27)	90.32 $\pm$ 0.31(95)(2)	96.67 $\pm$ 0.02(1)	97.16 $\pm$ 0.02(1)	98.75 $\pm$ 0.18(16)	96.52 $\pm$ 0.53(24)
cardio	96.59 $\pm$ 0.01(1)	96.59 $\pm$ 0.01(1)	88.94 $\pm$ 0.93(18)	80.71 $\pm$ 0.02(1)	69.21 $\pm$ 0.28(21)	87.86 $\pm$ 0.33(15)	56.34 $\pm$ 1.19(31)	99.46 $\pm$ 0.09(3)	99.14 $\pm$ 0.22(11)	96.96 $\pm$ 0.92(23)	98.13 $\pm$ 0.35(21)	89.53 $\pm$ 0.12(27)	90.32 $\pm$ 0.31(95)(2)	96.67 $\pm$ 0.02(1)	97.16 $\pm$ 0.02(1)	98.75 $\pm$ 0.18(16)	96.52 $\pm$ 0.53(24)
celebs	80.32 $\pm$ 0.04(1)	80.32 $\pm$ 0.04(1)	71.64 $\pm$ 0.91(18)	62.51 $\pm$ 0.02(1)	62.77 $\pm$ 0.82(16)	76.06 $\pm$ 1.44(4)	41.79 $\pm$ 0.02(1)	66.35 $\pm$ 0.03(1)	76.33 $\pm$ 0.63(11)	62.21 $\pm$ 1.88(27)	98.13 $\pm$ 0.35(21)	89.53 $\pm$ 0.12(27)	90.32 $\pm$ 0.31(95)(2)	96.67 $\pm$ 0.02(1)	97.16 $\pm$ 0.02(1)	98.75 $\pm$ 0.18(16)	96.52 $\pm$ 0.53(24)
ceus	70.53 $\pm$ 0.02(1)	70.53 $\pm$ 0.02(1)	59.33 $\pm$ 2.89(21)	62.51 $\pm$ 0.02(1)	62.77 $\pm$ 0.82(16)	76.06 $\pm$ 1.44(4)	41.79 $\pm$ 0.02(1)	66.35 $\pm$ 0.03(1)	76.33 $\pm$ 0.63(11)	62.21 $\pm$ 1.88(27)	98.13 $\pm$ 0.35(21)	89.53 $\pm$ 0.12(27)	90.32 $\pm$ 0.31(95)(2)	96.67 $\pm$ 0.02(1)	97.16 $\pm$ 0.02(1)	98.75 $\pm$ 0.18(16)	96.52 $\pm$ 0.53(24)
cover	94.33 $\pm$ 0.15(16)	94.33 $\pm$ 0.15(16)	76.35 $\pm$ 1.91(11)(24)	76.35 $\pm$ 1.91(11)(24)	76.35 $\pm$ 1.91(11)(24)	13.83 $\pm$ 1.34(32)	61.45 $\pm$ 0.93(1)	75.72 $\pm$ 0.59(12)	91.86 $\pm$ 0.03(5)	54.16 $\pm$ 0.33(27)	51.12 $\pm$ 1.19(29)	52.24 $\pm$ 1.19(29)	68.89 $\pm$ 1.11(22)	72.11 $\pm$ 0.16(4)	71.84 $\pm$ 0.16(4)	72.25 $\pm$ 1.62(5)	61.41 $\pm$ 0.71(9)
donors	88.63 $\pm$ 0.25(9)	88.63 $\pm$ 0.25(9)	91.64 $\pm$ 0.52(1)	81.09 $\pm$ 0.02(5)	75.34 $\pm$ 1.69(27)	76.35 $\pm$ 1.91(11)(24)	61.45 $\pm$ 0.93(1)	75.72 $\pm$ 0.59(12)	91.86 $\pm$ 0.03(5)	54.16 $\pm$ 0.33(27)	51.12 $\pm$ 1.19(29)	52.24 $\pm$ 1.19(29)	68.89 $\pm$ 1.11(22)	72.11 $\pm$ 0.16(4)	71.84 $\pm$ 0.16(4)	72.25 $\pm$ 1.62(5)	61.41 $\pm$ 0.71(9)
fruit	53.87 $\pm$ 0.00(25)	53.87 $\pm$ 0.00(25)	57.51 $\pm$ 1.51(7)	53.06 $\pm$ 0.02(6)	59.52 $\pm$ 1.44(8)	38.89 $\pm$ 0.01(4)	62.31 $\pm$ 0.02(5)	49.14 $\pm$ 0.02(6)	50.37 $\pm$ 0.02(8)	43.31 $\pm$ 1.02(25)	50.27 $\pm$ 1.88(29)	42.85 $\pm$ 1.92(7)	55.73 $\pm$ 0.33(25)	62.47 $\pm$ 0.02(1)	60.22 $\pm$ 0.07(7)	62.47 $\pm$ 0.02(1)	53.41 $\pm$ 0.63(24)
fruit	53.87 $\pm$ 0.00(25)	53.87 $\pm$ 0.00(25)	57.51 $\pm$ 1.51(7)	53.06 $\pm$ 0.02(6)	59.52 $\pm$ 1.44(8)	38.89 $\pm$ 0.01(4)	62.31 $\pm$ 0.02(5)	49.14 $\pm$ 0.02(6)	50.37 $\pm$ 0.02(8)	43.31 $\pm$ 1.02(25)	50.27 $\pm$ 1.88(29)	42.85 $\pm$ 1.92(7)	55.73 $\pm$ 0.33(25)	62.47 $\pm$ 0.02(1)	60.22 $\pm$ 0.07(7)	62.47 $\pm$ 0.02(1)	53.41 $\pm$ 0.63(24)
fruit	53.87 $\pm$ 0.00(25)	53.87 $\pm$ 0.00(25)	57.51 $\pm$ 1.51(7)	53.06 $\pm$ 0.02(6)	59.52 $\pm$ 1.44(8)	38.89 $\pm$ 0.01(4)	62.31 $\pm$ 0.02(5)	49.14 $\pm$ 0.02(6)	50.37 $\pm$ 0.02(8)	43.31 $\pm$ 1.02(25)	50.27 $\pm$ 1.88(29)	42.85 $\pm$ 1.92(7)	55.73 $\pm$ 0.33(25)	62.47 $\pm$ 0.02(1)	60.22 $\pm$ 0.07(7)	62.47 $\pm$ 0.02(1)	53.41 $\pm$ 0.63(24)
fruit	53.87 $\pm$ 0.00(25)	53.87 $\pm$ 0.00(25)	57.51 $\pm$ 1.51(7)	53.06 $\pm$ 0.02(6)	59.52 $\pm$ 1.44(8)	38.89 $\pm$ 0.01(4)	62.31 $\pm$ 0.02(5)	49.14 $\pm$ 0.02(6)	50.37 $\pm$ 0.02(8)	43.31 $\pm$ 1.02(25)	50.27 $\pm$ 1.88(29)	42.85 $\pm$ 1.92(7)	55.73 $\pm$ 0.33(25)	62.47 $\pm$ 0.02(1)	60.22 $\pm$ 0.07(7)	62.47 $\pm$ 0.02(1)	53.41 $\pm$ 0.63(24)
fruit	53.87 $\pm$ 0.00(25)	53.87 $\pm$ 0.00(25)	57.51 $\pm$ 1.51(7)	53.06 $\pm$ 0.02(6)	59.52 $\pm$ 1.44(8)	38.89 $\pm$ 0.01(4)	62.31 $\pm$ 0.02(5)	49.14 $\pm$ 0.02(6)	50.37 $\pm$ 0.02(8)	43.31 $\pm$ 1.02(25)	50.27 $\pm$ 1.88(29)	42.85 $\pm$ 1.92(7)	55.73 $\pm$ 0.33(25)	62.47 $\pm$ 0.02(1)	60.22 $\pm$ 0.07(7)	62.47 $\pm$ 0.02(1)	53.41 $\pm$ 0.63(24)
fruit	53.87 $\pm$ 0.00(25)	53.87 $\pm$ 0.00(25)	57.51 $\pm$ 1.51(7)	53.06 $\pm$ 0.02(6)	59.52 $\pm$ 1.44(8)	38.89 $\pm$ 0.01(4)	62.31 $\pm$ 0.02(5)	49.14 $\pm$ 0.02(6)	50.37 $\pm$ 0.02(8)	43.31 $\pm$ 1.02(25)	50.27 $\pm$ 1.88(29)	42.85 $\pm$ 1.92(7)	55.73 $\pm$ 0.33(25)	62.47 $\pm$ 0.02(1)	60.22 $\pm$ 0.07(7)	62.47 $\pm$ 0.02(1)	53.41 $\pm$ 0.63(24)
fruit	53.87 $\pm$ 0.00(25)	53.87 $\pm$ 0.00(25)	57.51 $\pm$ 1.51(7)	53.06 $\pm$ 0.02(6)	59.52 $\pm$ 1.44(8)	38.89 $\pm$ 0.01(4)	62.31 $\pm$ 0.02(5)	49.14 $\pm$ 0.02(6)	50.37 $\pm$ 0.02(8)	43.31 $\pm$ 1.02(25)	50.27 $\pm$ 1.88(29)	42.85 $\pm$ 1.92(7)	55.73 $\pm$ 0.33(25)	62.47 $\pm$ 0.02(1)	60.22 $\pm$ 0.07(7)	62.47 $\pm$ 0.02(1)	53.41 $\pm$ 0.63(24)
fruit	53.87 $\pm$ 0.00(25)	53.87 $\pm$ 0.00(25)	57.51 $\pm$ 1.51(7)	53.06 $\pm$ 0.02(6)	59.52 $\pm$ 1.44(8)	38.89 $\pm$ 0.01(4)	62.31 $\pm$ 0.02(5)	49.14 $\pm$ 0.02(6)	50.37 $\pm$ 0.02(8)	43.31 $\pm$ 1.02(25)	50.27 $\pm$ 1.88(29)	42.85 $\pm$ 1.92(7)	55.73 $\pm$ 0.33(25)	62.47 $\pm$ 0.02(1)	60.22 $\pm$ 0.07(7)	62.47 $\pm$ 0.02(1)	53.41 $\pm$ 0.63(24)
fruit	53.87 $\pm$ 0.00(25)	53.87 $\pm$ 0.00(25)	57.51 $\pm$ 1.51(7)	53.06 $\pm$ 0.02(6)	59.52 $\pm$ 1.44(8)	38.89 $\pm$ 0.01(4)	62.31 $\pm$ 0.02(5)	49.14 $\pm$ 0.02(6)	50.37 $\pm$ 0.02(8)	43.31 $\pm$ 1.02(25)	50.27 $\pm$ 1.88(29)	42.85 $\pm$ 1.92(7)	55.73 $\pm$ 0.33(25)	62.47 $\pm$ 0.02(1)	60.22 $\pm$ 0.07(7)	62.47 $\pm$ 0.02(1)	53.41 $\pm$ 0.63(24)
fruit	53.87 $\pm$ 0.00(25)	53.87 $\pm$ 0.00(25)	57.51 $\pm$ 1.51(7)	53.06 $\pm$ 0.02(6)	59.52 $\pm$ 1.44(8)	38.89 $\pm$ 0.01(4)	62.31 $\pm$ 0.02(5)	49.14 $\pm$ 0.02(6)	50.37 $\pm$ 0.02(8)	43.31 $\pm$ 1.02(25)	50.27 $\pm$ 1.88(29)	42.85 $\pm$ 1.92(7)	55.73 $\pm$ 0.33(25)	62.47 $\pm$ 0.02(1)	60.22 $\pm$ 0.07(7)	62.47 $\pm$ 0.02(1)	53.41 $\pm$ 0.63(24)
fruit	53.87 $\pm$ 0.00(25)	53.87 $\pm$ 0.00(25)	57.51 $\pm$ 1.51(7)	53.06 $\pm$ 0.02(6)	59.52 $\pm$ 1.44(8)	38.89 $\pm$ 0.01(4)	62.31 $\pm$ 0.02(5)	49.14 $\pm$ 0.02(6)	50.37 $\pm$ 0.02(8)	43.31 $\pm$ 1.02(25)	50.27 $\pm$ 1.88(29)	42.85 $\pm$ 1.92(7)	55.73 $\pm$ 0.33(25)	62.47 $\pm$ 0.02(1)	60.22 $\pm$ 0.07(7)	62.47 $\pm$ 0.02(1)	53.41 $\pm$ 0.63(24)
fruit	53.87 $\pm$ 0.00(25)	53.87 $\pm$ 0.00(25)	57.51 $\pm$ 1.51(7)	53.06 $\pm$ 0.02(6)	59.52 $\pm$ 1.44(8)	38.89 $\pm$ 0.01(4)	62.31 $\pm$ 0.02(5)	49.14 $\pm$ 0.02(6)	50.37 $\pm$ 0.02(8)	43.31 $\pm$ 1.02(25)	50.27 $\pm$ 1.88(29)	42.85 $\pm$ 1.92(7)	55.73 $\pm$ 0.33(25)	62.47 $\pm$ 0.02(1)	60.22 $\pm$ 0.07(7)	62.47 $\pm$ 0.02(1)	53.41 $\pm$ 0.63(24)
fruit	53.87 $\pm$ 0.00(25)	53.87 $\pm$ 0.00(25)	57.51 $\pm$ 1.51(7)	53.06 $\pm$ 0.02(6)	59.52 $\pm$ 1.44(8)	38.89 $\pm$ 0.01(4)	62.31 $\pm$ 0.02(5)	49.14 $\pm$ 0.02(6)	50.37 $\pm$ 0.02(8)	43.31 $\pm$ 1.02(25)	50.27 $\pm$ 1.88(29)	42.85 $\pm$ 1.92(7)	55.73 $\pm$ 0.33(25)	62.47 $\pm$ 0.02(1)	60.22 $\pm$ 0.07(7)	62.47 $\pm$ 0.02(1)	53.41 $\pm$ 0.63(24)
fruit	53.87 $\pm$ 0.00(25)	53.87 $\pm$ 0.00(25)	57.51 $\pm$ 1.51(7)	53.06 $\pm$ 0.02(6)	59.52 $\pm$ 1.44(8)	38.89 $\pm$ 0.01(4)	62.31 $\pm$ 0.02(5)	49.14 $\pm$ 0.02(6)	50.37 $\pm$ 0.02(8)	43.31 $\pm$ 1.02(25)	50.27 $\pm$ 1.88(29)	42.85 $\pm$ 1.92(7)	55.73 $\pm$ 0.33(25)	62.47 $\pm$ 0.02(1)	60.22 $\pm$ 0.07(7)	62.47 $\pm$ 0.02(1)	53.41 $\pm$ 0.63(24)
fruit	53.87 $\pm$ 0.00(25)	53.87 $\pm$ 0.00(25)	57.51 $\pm$ 1.51(7)	53.06 $\pm$ 0.02(6)	59.52 $\pm$ 1.44(8)	38.89 $\pm$ 0.01(4)	62.31 $\pm$ 0.02(5)	49.14 $\pm$ 0.02(6)	50.37 $\pm$ 0.02(8)	43.31 $\pm$ 1.02(25)	50.27 $\pm$ 1.88(29)	42.85 $\pm$ 1.92(7)	55.73 $\pm$ 0.33(25)	62.47 $\pm$ 0.02(1)	60.22 $\pm$ 0.07(7)	62.47 $\pm$ 0.02(1)	53.41 $\pm$ 0.63(24)
fruit	53.87 $\pm$ 0.00(25)	53.87 $\pm$ 0.00(25)	57.51 $\pm$ 1.51(7)	53.06 $\pm$ 0.02(6)	59.52 $\pm$ 1.44(8)	38.89 $\pm$ 0.01(4)	62.31 $\pm$ 0.02(5)	49.14 $\pm$ 0.02(6)	50.37 $\pm$ 0.02(8)	43.31 $\pm$ 1.02(25)	50.27 $\pm$ 1.88(29)	42.85 $\pm$ 1.92(7)	55.73 $\pm$ 0.33(25)	62.47 $\pm$ 0.02(1)	60.22 $\pm$ 0.07(7)	62.47 $\pm$ 0.02(1)	53.41 $\pm$ 0.63(24)
fruit	53.87 $\pm$ 0.00(25)	53.87 $\pm$ 0.00(25)	57.51 $\pm$ 1.51(7)	53.06 $\pm$ 0.02(6)	59.52 $\pm$ 1.44(8)	38.89 $\pm$ 0.01(4)	62.31 $\pm$ 0.02(5)	49.14 $\pm$ 0.02(6)	50.37 $\pm$ 0.02(8)	43.31 $\pm$ 1.02(25)	50.27 $\pm$ 1.88(29)	42.85 $\pm$ 1.92(7)	55.73 $\pm$ 0.33(25)	62.47 $\pm$ 0.02(1)	60.22 $\pm$ 0.07(7)	62.47 $\pm$ 0.02(1)	53.41 $\pm$ 0.63(24)
fruit	53.87 $\pm$ 0.00(25)	53.87 $\pm$ 0.00(25)	57.51 $\pm$ 1.51(7)	53.06 $\pm$ 0.02(6)	59.52 $\pm$ 1.44(8)	38.89 $\pm$ 0.01(4)	62.31 $\pm$ 0.02(5)	49.14 $\pm$ 0.02(6)	50.37 $\pm$ 0.02(8)	43.31 $\pm$ 1.02(25)	50.27 $\pm$ 1.88(29)	42.85 $\pm$ 1.92(7)	55.73 $\pm$ 0.33(25)	62.47 $\pm$ 0.02(1)	60.22 $\pm$ 0.07(7)	62.47 $\pm$ 0.02(1)	53.41 $\pm$ 0.63(24)
fruit	53.87 $\pm$ 0.00(25)	53.87 $\pm$ 0.00(25)	57.51 $\pm$ 1.51(7)	53.06 $\pm$ 0.02(6)	59.52 $\pm$ 1.44(8)	38.89 $\pm$ 0.01(4)	62.31 $\pm$ 0.02(5)	49.14 $\pm$ 0.02(6)	50.37 $\pm$ 0.02(8)	43.31 $\pm$ 1.02(25)	50.27 $\pm$ 1.88(29)	42.85 $\pm$ 1.92(7)	55.73 $\pm$ 0.33(25)	62.47 $\pm$ 0.02(1)	60.22 $\pm$ 0.07(7)	62.47 $\pm$ 0.02(1)	53.41 $\pm$ 0.63(24)
fruit	53.87 $\pm$ 0.00(25)	53.87 $\pm$ 0.00(25)	57.51 $\pm$ 1.51(7)	53.06 $\pm$ 0.02(6)	59.52 $\pm$ 1.44(8)	38.89 $\pm$ 0.01(4)	62.31 $\pm$ 0.02(5)	49.14 $\pm$ 0.02(6)	50.37 $\pm$ 0.02(8)	43.31 $\pm$ 1.02(25)	50.27 $\pm$ 1.88(29)	42.85 $\pm$ 1.92(7)	55.73 $\pm$ 0.33(25)	62.47 $\pm$ 0.02(1)	60.22 $\pm$ 0.07(7)	62.47 $\pm$ 0.02(1)	53.41 $\pm$ 0.63(24)
fruit	53.87 $\pm$ 0.00(25)	53.87 $\pm$ 0.00(25)	57.51 $\pm$ 1.51(														

Table 17.1: 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.

Dataset	F0+D0 (D=100)	F0+D0 (D=250)	DTE-NN	ANN	KL	DTEC	LOF	CHLOF	Feature-Ranking	SLAD	DDPM	QCSVM	DTE-AG	Preest	MCD
amazon	6.93E-04(5.05)	6.02E-04(0.65)	6.02E-04(0.65)	6.02E-04(0.65)	5.53E-04(3.30)	5.76E-04(2.26)	6.54E-04(5)	6.43E-04(9)	6.63E-04(9)	5.99E-04(8.18)	5.97E-04(8.20)	6.32E-04(0.7)	5.98E-04(1.4)	5.52E-04(1.2)	5.55E-04(0.29)
	1.1E-04(0.52)	1.09E-04(0.5)	1.09E-04(0.5)	1.09E-04(0.5)	1.04E-04(0.4)	1.04E-04(0.6)	1.04E-04(0.6)	1.04E-04(0.6)	1.06E-04(0.92)	1.06E-04(0.92)	1.06E-04(0.92)	1.10E-04(0.12)	1.06E-04(0.12)	1.07E-04(0.11)	1.07E-04(0.11)
	4.38E-04(1.04)	4.73E-04(1.00)	4.73E-04(1.00)	4.73E-04(1.00)	4.64E-04(1.1)	4.64E-04(1.1)	5.57E-04(2.7)	5.57E-04(2.7)	5.57E-04(2.7)	4.83E-04(1.2)	4.83E-04(1.2)	4.83E-04(1.2)	8.19E-04(2.5)	9.72E-04(2.1)	22.16E-03(1.7)
breast	9.91E-04(0.41)	9.89E-04(0.41)	9.89E-04(0.41)	9.89E-04(0.41)	9.63E-04(1.1)	9.63E-04(1.1)	9.09E-04(2.41)	9.09E-04(2.41)	9.29E-04(0.98)	9.29E-04(0.98)	9.35E-04(0.2)	9.35E-04(0.2)	8.41E-04(2.628)	9.63E-04(0.21)	9.63E-04(0.21)
	34.24E-04(7.26)	37.78E-04(1.02)	37.78E-04(1.02)	37.78E-04(1.02)	40.99E-04(0.89)	40.99E-04(0.89)	40.24E-04(0.21)	40.24E-04(0.21)	38.56E-04(0.91)	38.56E-04(0.91)	48.97E-04(2.06)	48.97E-04(2.06)	46.73E-04(1.85)	47.91E-04(1.96)	47.91E-04(1.96)
	78.92E-04(0.01)	76.73E-04(0.01)	76.73E-04(0.01)	76.73E-04(0.01)	77.05E-04(0.01)	77.05E-04(0.01)	75.53E-04(0.89)	75.53E-04(0.89)	71.55E-04(1.34)	71.55E-04(1.34)	69.94E-04(0.20)	69.94E-04(0.20)	41.07E-04(1.30)	78.63E-04(0.20)	78.63E-04(0.20)
cardiography	6.35E-04(0.00)	5.74E-04(0.00)	5.74E-04(0.00)	5.74E-04(0.00)	5.73E-04(0.00)	5.73E-04(0.00)	6.17E-04(0.28)	6.17E-04(0.28)	5.94E-04(0.92)	5.94E-04(0.92)	6.15E-04(0.92)	6.15E-04(0.92)	39.99E-04(0.33)	62.85E-04(0.27)	62.85E-04(0.27)
	1.60E-04(2.71)	1.60E-04(2.71)	1.60E-04(2.71)	1.60E-04(2.71)	1.74E-04(1.03)	1.74E-04(1.03)	2.03E-04(2.26)	2.03E-04(2.26)	1.92E-04(2.38)	1.92E-04(2.38)	1.92E-04(2.38)	2.23E-04(0.69)	1.63E-04(1.75)	1.49E-04(2.72)	2.89E-04(1.47)
	72.57E-04(8.86)	68.27E-04(2.59)	68.27E-04(2.59)	68.27E-04(2.59)	34.88E-04(6.16)	34.88E-04(6.16)	73.33E-04(3.37)	73.33E-04(3.37)	69.95E-04(6.28)	69.95E-04(6.28)	73.73E-04(3.4)	22.8E-04(1.00)	8.66E-04(1.52)	1.14E-04(1.828)	1.14E-04(1.828)
cones	96.43E-04(8.48)	95.55E-04(8.48)	95.55E-04(8.48)	95.55E-04(8.48)	98.35E-04(0.97)	98.35E-04(0.97)	63.39E-04(1.81)	63.39E-04(1.81)	65.35E-04(1.20)	65.35E-04(1.20)	26.66E-04(3.28)	42.1E-04(0.86)	9.73E-04(2.6)	40.51E-04(3.90)	31.24E-04(6.26)
	63.08E-04(0.08)	60.86E-04(3.58)	60.86E-04(3.58)	60.86E-04(3.58)	63.18E-04(7.0)	63.18E-04(7.0)	50.44E-04(0.32)	50.44E-04(0.32)	61.29E-04(1.85)	61.29E-04(1.85)	64.75E-04(0.62)	61.2E-04(0.61)	63.83E-04(1.25)	59.22E-04(0.22)	60.06E-04(3.96)
	55.13E-04(1.78)	34.75E-04(78.42)	34.75E-04(78.42)	34.75E-04(78.42)	33.88E-04(7.0)	33.88E-04(7.0)	55.69E-04(0.19)	55.69E-04(0.19)	27.77E-04(1.427)	27.77E-04(1.427)	69.31E-04(6.38)	54.97E-04(0.04)	51.4E-04(6.61)	19.22E-04(0.26)	60.06E-04(3.96)
fruit	8.56E-04(4.52)	70.36E-04(9.67)	70.36E-04(9.67)	70.36E-04(9.67)	92.35E-04(8.21)	92.35E-04(8.21)	38.12E-04(5.91)	38.12E-04(5.91)	31.62E-04(4.7)	31.62E-04(4.7)	41.55E-04(3.98)	27.66E-04(3.98)	80.57E-04(12.83)	21.37E-04(3.98)	25.39E-04(2.227)
	99.48E-04(1.4)	95.82E-04(0.95)	95.82E-04(0.95)	95.82E-04(0.95)	91.51E-04(3.68)	91.51E-04(3.68)	43.67E-04(0.95)	43.67E-04(0.95)	33.74E-04(8.25)	33.74E-04(8.25)	39.11E-04(1.49)	39.11E-04(1.49)	9.95E-04(0.25)	55.36E-04(0.28)	58.82E-04(0.2)
	99.48E-04(1.4)	95.82E-04(0.95)	95.82E-04(0.95)	95.82E-04(0.95)	91.51E-04(3.68)	91.51E-04(3.68)	43.67E-04(0.95)	43.67E-04(0.95)	33.74E-04(8.25)	33.74E-04(8.25)	39.11E-04(1.49)	39.11E-04(1.49)	9.95E-04(0.25)	55.36E-04(0.28)	58.82E-04(0.2)
hip	1.60E-04(2.71)	1.60E-04(2.71)	1.60E-04(2.71)	1.60E-04(2.71)	1.74E-04(1.03)	1.74E-04(1.03)	2.03E-04(2.26)	2.03E-04(2.26)	1.92E-04(2.38)	1.92E-04(2.38)	1.92E-04(2.38)	2.23E-04(0.69)	1.63E-04(1.75)	1.49E-04(2.72)	2.89E-04(1.47)
	72.57E-04(8.86)	68.27E-04(2.59)	68.27E-04(2.59)	68.27E-04(2.59)	34.88E-04(6.16)	34.88E-04(6.16)	73.33E-04(3.37)	73.33E-04(3.37)	69.95E-04(6.28)	69.95E-04(6.28)	73.73E-04(3.4)	22.8E-04(1.00)	8.66E-04(1.52)	1.14E-04(1.828)	1.14E-04(1.828)
	96.43E-04(8.48)	95.55E-04(8.48)	95.55E-04(8.48)	95.55E-04(8.48)	98.35E-04(0.97)	98.35E-04(0.97)	63.39E-04(1.81)	63.39E-04(1.81)	65.35E-04(1.20)	65.35E-04(1.20)	26.66E-04(3.28)	42.1E-04(0.86)	9.73E-04(2.6)	40.51E-04(3.90)	31.24E-04(6.26)
denon	63.08E-04(0.08)	60.86E-04(3.58)	60.86E-04(3.58)	60.86E-04(3.58)	63.18E-04(7.0)	63.18E-04(7.0)	50.44E-04(0.32)	50.44E-04(0.32)	61.29E-04(1.85)	61.29E-04(1.85)	64.75E-04(0.62)	61.2E-04(0.61)	63.83E-04(1.25)	59.22E-04(0.22)	60.06E-04(3.96)
	55.13E-04(1.78)	34.75E-04(78.42)	34.75E-04(78.42)	34.75E-04(78.42)	33.88E-04(7.0)	33.88E-04(7.0)	55.69E-04(0.19)	55.69E-04(0.19)	27.77E-04(1.427)	27.77E-04(1.427)	69.31E-04(6.38)	54.97E-04(0.04)	51.4E-04(6.61)	19.22E-04(0.26)	60.06E-04(3.96)
	8.56E-04(4.52)	70.36E-04(9.67)	70.36E-04(9.67)	70.36E-04(9.67)	92.35E-04(8.21)	92.35E-04(8.21)	38.12E-04(5.91)	38.12E-04(5.91)	31.62E-04(4.7)	31.62E-04(4.7)	41.55E-04(3.98)	27.66E-04(3.98)	80.57E-04(12.83)	21.37E-04(3.98)	25.39E-04(2.227)
legatus	99.48E-04(1.4)	95.82E-04(0.95)	95.82E-04(0.95)	95.82E-04(0.95)	91.51E-04(3.68)	91.51E-04(3.68)	43.67E-04(0.95)	43.67E-04(0.95)	33.74E-04(8.25)	33.74E-04(8.25)	39.11E-04(1.49)	39.11E-04(1.49)	9.95E-04(0.25)	55.36E-04(0.28)	58.82E-04(0.2)
	99.48E-04(1.4)	95.82E-04(0.95)	95.82E-04(0.95)	95.82E-04(0.95)	91.51E-04(3.68)	91.51E-04(3.68)	43.67E-04(0.95)	43.67E-04(0.95)	33.74E-04(8.25)	33.74E-04(8.25)	39.11E-04(1.49)	39.11E-04(1.49)	9.95E-04(0.25)	55.36E-04(0.28)	58.82E-04(0.2)
	99.48E-04(1.4)	95.82E-04(0.95)	95.82E-04(0.95)	95.82E-04(0.95)	91.51E-04(3.68)	91.51E-04(3.68)	43.67E-04(0.95)	43.67E-04(0.95)	33.74E-04(8.25)	33.74E-04(8.25)	39.11E-04(1.49)	39.11E-04(1.49)	9.95E-04(0.25)	55.36E-04(0.28)	58.82E-04(0.2)
hip	1.60E-04(2.71)	1.60E-04(2.71)	1.60E-04(2.71)	1.60E-04(2.71)	1.74E-04(1.03)	1.74E-04(1.03)	2.03E-04(2.26)	2.03E-04(2.26)	1.92E-04(2.38)	1.92E-04(2.38)	1.92E-04(2.38)	2.23E-04(0.69)	1.63E-04(1.75)	1.49E-04(2.72)	2.89E-04(1.47)
	72.57E-04(8.86)	68.27E-04(2.59)	68.27E-04(2.59)	68.27E-04(2.59)	34.88E-04(6.16)	34.88E-04(6.16)	73.33E-04(3.37)	73.33E-04(3.37)	69.95E-04(6.28)	69.95E-04(6.28)	73.73E-04(3.4)	22.8E-04(1.00)	8.66E-04(1.52)	1.14E-04(1.828)	1.14E-04(1.828)
	96.43E-04(8.48)	95.55E-04(8.48)	95.55E-04(8.48)	95.55E-04(8.48)	98.35E-04(0.97)	98.35E-04(0.97)	63.39E-04(1.81)	63.39E-04(1.81)	65.35E-04(1.20)	65.35E-04(1.20)	26.66E-04(3.28)	42.1E-04(0.86)	9.73E-04(2.6)	40.51E-04(3.90)	31.24E-04(6.26)
internals	36.29E-04(2.227)	30.65E-04(1.25)	31.32E-04(2.00)	40.92E-04(0.03)	40.03E-04(1.364)	55.22E-04(3.67)	50.04E-04(0.01)	47.80E-04(0.02)	49.26E-04(1.812)	49.26E-04(1.812)	47.72E-04(0.16)	88.5E-04(0.015)	58.6E-04(0.225)	89.24E-04(1.742)	34.56E-04(0.28)
	36.29E-04(2.227)	30.65E-04(1.25)	31.32E-04(2.00)	40.92E-04(0.03)	40.03E-04(1.364)	55.22E-04(3.67)	50.04E-04(0.01)	47.80E-04(0.02)	49.26E-04(1.812)	49.26E-04(1.812)	47.72E-04(0.16)	88.5E-04(0.015)	58.6E-04(0.225)	89.24E-04(1.742)	34.56E-04(0.28)
	36.29E-04(2.227)	30.65E-04(1.25)	31.32E-04(2.00)	40.92E-04(0.03)	40.03E-04(1.364)	55.22E-04(3.67)	50.04E-04(0.01)	47.80E-04(0.02)	49.26E-04(1.812)	49.26E-04(1.812)	47.72E-04(0.16)	88.5E-04(0.015)	58.6E-04(0.225)	89.24E-04(1.742)	34.56E-04(0.28)
isorepine	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)
	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)
	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)	97.75E-04(0.768)
landuse	58.24E-04(0.02)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)
	58.24E-04(0.02)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)
	58.24E-04(0.02)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)	54.75E-04(23.47)
letter	9.33E-04(1.02)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)
	9.33E-04(1.02)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)
	9.33E-04(1.02)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)	9.13E-04(39.01)
lymphogly	36.29E-04(2.227)	30.65E-04(1.25)	31.32E-04(2.00)	40.92E-04(0.03)	40.03E-04(1.364)	55.22E-04(3.67)	50.04E-04(0.01)	47.80E-04(0.02)	49.26E-04(1.812)	49.26E-04(1.812)	47.72E-04(0.16)	88.5E-04(0.015)	58.6E-04(0.225)	89.24E-04(1.742)	34.56E-04(0.28)
	36.29E-04(2.227)	30.65E-04(1.25)	31.32E-04(2.00)	40.92E-04(0.03)	40.03E-04(1.364)	55.22E-04(3.67)	50.04E-04(0.01)	47.80E-04(0.02)	49.26E-04(1.812)	49.26E-04(1.812)	47.72E-04(0.16)	88.5E-04(0.015)	58.6E-04(0.225)	89.24E-04(1.742)	34.56E-04(0.28)
	36.29E-04(2.227)	30.65E-04(1.25)	31.32E-04(2.00)	40.92E-04(0.03)	40.03E-04(1.364)	55.22E-04(3.67)	50.04E-04(0.01)	47.80E-04(0.02)	49.26E-04(1.812)	49.26E-04(1.812)	47.72E-04(0.16)	88.5E-04(0.015)	58.6E-04(0.225)	89.24E-04(1.742)	34.56E-04(0.28)
morphology	97.75E-04(0.015)	97.75E-04(0													

Table 17.2: 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.

Dataset	VAE	PCA	PlanarFlow	HBOs	GANomaly	QOAD	DIF	COPOD	ECOD	DeepSVDD	LODA	DAGMM	DBOCC	DTE-NP ⁺⁺	KN ⁺⁺	ICL ⁺⁺	DTE-C ⁺⁺
aldi	6.54±0.00(5)	6.54±0.00(5)	5.48±0.00(2)	6.42±0.00(8)	8.09±1.27(1)	5.7±0.24(28)	5.8±0.00(2)	5.72±0.00(2)	6.06±0.00(4)	6.23±0.38(0.5)	5.93±0.52(1)	6.07±0.55(13)	5.91±0.00(2)	6.02±0.00(16.5)	6.09±0.00(2)	5.49±0.02(31)	5.79±0.00(25)
amazon	10.72±0.00(7.5)	10.72±0.00(7.5)	9.56±0.59(30)	11.1±0.00(9.5)	9.93±0.10(21)	10.94±0.21(15)	9.89±0.41(28)	11.15±0.00(7.5)	10.40±0.00(2)	10.24±1.27(22)	10.18±0.78(24)	9.53±0.17(28)	9.53±0.17(28)	10.61±0.00(4.5)	10.52±0.00(5)	10.01±0.00(26)	10.42±0.40(20)
anthyroid	56.74±0.00(17)	56.75±0.00(17)	65.15±8.65(8)	39.03±0.00(2)	34.3±0.62(30)	58.75±0.14(2)	61.12±0.00(3)	29.01±0.00(28)	40.02±0.00(28)	27.83±5.95(2)	49.01±0.76(23)	48.03±1.36(26)	63.72±3.08(9)	66.11±0.00(7)	67.06±0.00(5)	52.44±0.00(23)	68.86±0.2(3)
backdoor	7.97±0.02(424)	7.94±0.12(25)	32.15±2.78(14)	8.56±0.26(23)	27.74±6.65(26)	6.31±1.94(28)	1.78±1.10(318)	4.84±0.11(8.5)	4.84±0.11(8.5)	84.77±2.79(3)	5.96±1.38(29)	7.53±1.45(27)	7.53±1.45(27)	40.48±0.00(12)	39.08±0.23(5)	87.63 ±1.05(2)	35.68±0.79(15)
breast	99.12±0.00(1.5)	99.12±0.00(1.5)	97.12±0.00(1.5)	99.08±0.02(9.5)	97.78±0.00(1.5)	98.77±0.00(1.5)	98.77±0.00(1.5)	99.08±0.02(9.5)	99.16±0.00(3)	96.51±1.28(23)	96.51±1.28(23)	90.93±1.73(26)	63.19±2.13(54)	97.69±0.00(19)	97.69±0.00(19)	97.69±0.00(19)	94.04±0.69(24)
campaign	40.02±0.00(1.5)	40.02±0.00(1.5)	42.77±0.29(20)	42.77±0.29(20)	39.17±4.25(22)	23.69±1.78(31)	24.69±0.00(30)	31.18 ±0.00(1)	49.31±0.00(3)	38.95±1.27(25)	32.53±2.92(29)	32.53±2.92(29)	20.25±0.00(52)	49.31±0.00(7)	49.31±0.00(7)	40.02±0.00(1.5)	37.11±0.64(26)
catdog	40.02±0.00(1.5)	40.02±0.00(1.5)	42.77±0.29(20)	42.77±0.29(20)	39.17±4.25(22)	23.69±1.78(31)	24.69±0.00(30)	31.18 ±0.00(1)	49.31±0.00(3)	38.95±1.27(25)	32.53±2.92(29)	32.53±2.92(29)	20.25±0.00(52)	49.31±0.00(7)	49.31±0.00(7)	40.02±0.00(1.5)	37.11±0.64(26)
celeba	19.82±0.11(5)	19.82±0.11(5)	16.77±0.78(8)	16.77±0.78(8)	17.47±0.21(14)	4.01±1.21(29)	7.09±1.06(24)	16.84±0.82(9)	7.09±1.06(24)	48.75±1.27(2)	48.75±1.27(2)	48.75±1.27(2)	48.75±1.27(2)	19.82±0.11(5)	19.82±0.11(5)	19.82±0.11(5)	19.82±0.11(5)
celeba-coverage	20.95±1.11(15)	20.95±1.11(15)	14.69±1.56(20)	14.69±1.56(20)	17.47±0.21(14)	4.01±1.21(29)	7.09±1.06(24)	16.84±0.82(9)	7.09±1.06(24)	48.75±1.27(2)	48.75±1.27(2)	48.75±1.27(2)	48.75±1.27(2)	20.95±1.11(15)	20.95±1.11(15)	20.95±1.11(15)	20.95±1.11(15)
cover	10.05±0.88(21)	10.05±0.88(21)	1.98±0.06(31)	5.42±0.00(1)	25.03±3.82(46)	1.09±1.93(32)	2.23±0.14(21)	12.26±0.88(23)	19.22±1.50(9)	15.35±1.07(17)	13.42±0.92(27)	13.74±0.92(28)	13.74±0.92(28)	12.63±0.51(15)	14.11±0.61(9)	13.74±0.92(28)	10.05±0.88(21)
donors	36.01±0.72(23)	35.22±1.21(24)	49.31±1.82(13)	36.33±1.85(22)	23.91±1.25(30)	62.14±5.52(10)	60.8±0.00(1)	33.51±0.80(25)	41.27±0.00(29)	22.56±9.15(29)	22.56±9.15(29)	9.84±0.98(24)	31.25±1.58(15)	51.95±2.00(9)	79.92±1.13(7)	76.75±3.27(8)	43.3±2.85(16)
fault	60.35±0.00(20)	60.35±0.00(20)	60.35±0.00(20)	60.35±0.00(20)	62.47±4.52(10)	62.14±5.52(10)	60.8±0.00(1)	33.51±0.80(25)	41.27±0.00(29)	55.46±1.48(26)	54.49±2.63(27)	56.75±1.66(25)	57.81±4.16(24)	61.84±0.00(4)	62.56±0.00(9)	64.4±1.09(3)	58.79±0.41(23)
frail	28.74±5.54(26)	26.93±1.19(128)	62.81±9.37(3)	32.25±5.42(22)	60.24±1.15(5)	29.44±2.46(25)	2.08±0.82(31)	38.43±1.99(18)	35.32±4.68(21)	48.33±1.71(31)	36.59±1.51(17)	55.71±0.11(30)	43.12±0.00(32)	41.22±0.50(15)	40.45±4.07(16)	30.33±4.24(25)	30.33±4.24(25)
glass	18.51±3.73(30)	20.96±5.76(26)	30.93±6.56(18)	27.61±6.50(22)	26.04±0.10(32)	18.33±7.43(9)	67.02±1.31(6)	20.09±4.04(28)	25.35±2.04(7)	52.35±1.07(4)	52.35±1.07(4)	18.62±1.42(39)	34.91±1.39(24)	41.38±4.46(10)	31.91±3.73(5)	76.88±1.33(84)	27.78±2.62(19)
hepatitis	64.48±5.12(1)	64.85±5.12(1)	89.63±5.09(22)	63.49±5.99(22)	73.25±1.71(517)	67.77±5.43(19)	89.63±5.12(1)	56.08±1.35(25)	55.84±5.47(29)	90.15±1.08(7)	90.15±1.08(7)	54.73±0.80(27)	34.91±1.39(24)	89.63±5.12(1)	89.63±5.12(1)	74.88±1.79(16)	74.88±1.79(16)
hnp	90.42±0.00(1.5)	91.69±1.53(10)	52.23±4.21(22)	52.23±4.21(22)	38.95±1.27(25)	23.69±1.78(31)	24.69±0.00(30)	31.18±0.00(1)	49.31±0.00(3)	38.95±1.27(25)	32.53±2.92(29)	32.53±2.92(29)	20.25±0.00(52)	90.42±0.00(1.5)	90.42±0.00(1.5)	43.08±2.22(4)	43.08±2.22(4)
imdb	40.02±0.00(1.5)	40.02±0.00(1.5)	42.77±0.29(20)	42.77±0.29(20)	39.17±4.25(22)	23.69±1.78(31)	24.69±0.00(30)	31.18 ±0.00(1)	49.31±0.00(3)	38.95±1.27(25)	32.53±2.92(29)	32.53±2.92(29)	20.25±0.00(52)	40.02±0.00(1.5)	40.02±0.00(1.5)	40.02±0.00(1.5)	37.11±0.64(26)
imdb-coverage	40.02±0.00(1.5)	40.02±0.00(1.5)	42.77±0.29(20)	42.77±0.29(20)	39.17±4.25(22)	23.69±1.78(31)	24.69±0.00(30)	31.18 ±0.00(1)	49.31±0.00(3)	38.95±1.27(25)	32.53±2.92(29)	32.53±2.92(29)	20.25±0.00(52)	40.02±0.00(1.5)	40.02±0.00(1.5)	40.02±0.00(1.5)	37.11±0.64(26)
landsat	91.42±1.37(24)	90.94±1.25(25)	97.14±0.94(27)	64.63±7.62(32)	95.38±1.33(17.5)	91.17±2.61(28)	94.44±1.96(21)	78.89±1.06(28)	31.09±0.32(1)	88.15±3.21(6)	85.15±3.72(24)	77.53±4.99(29)	71.72±1.21(883)	97.96±0.00(4)	96.01±0.78(16)	97.92±1.11(6)	86.53±4.28(26)
landsat-coverage	91.42±1.37(24)	90.94±1.25(25)	97.14±0.94(27)	64.63±7.62(32)	95.38±1.33(17.5)	91.17±2.61(28)	94.44±1.96(21)	78.89±1.06(28)	31.09±0.32(1)	88.15±3.21(6)	85.15±3.72(24)	77.53±4.99(29)	71.72±1.21(883)	97.96±0.00(4)	96.01±0.78(16)	97.92±1.11(6)	86.53±4.28(26)
letter	8.01±0.00(32)	8.01±0.00(32)	9.19±0.00(13)	8.73±0.00(18)	8.45±0.09(23)	8.13±0.05(28)	94.2 ±0.00(1)	8.85±0.07(1)	9.05±0.21(16.5)	8.93±0.52(16.5)	8.03±0.25(30)	10.37±1.67(8)	15.74 ±1.69(12)	8.07±0.00(21)	8.47±0.00(24)	11.41±0.00(2)	10.09±0.04(10)
lymphography	98.59±1.01(12)	98.49±0.55(14)	96.24±4.22(19)	96.55±2.11(18)	90.53±7.43(24)	98.76±2.44(2)	98.12±6.61(6)	93.92±4.85(23)	94.38±1.18(21.5)	96.82±3.09(17)	24.13±1.29(32)	73.47±1.88(28)	30.87±5.63(31)	96.16±6.43(20)	98.57±0.65(13)	100.0 ±0.00(5)	68.82±5.96(30)
magic-gamma	75.27±0.00(25)	75.27±0.00(25)	78.51±2.91(18)	77.15±0.00(2)	65.83±2.36(30)	76.13±2.42(23)	65.76±0.00(3)	72.22±0.00(27)	67.92±0.00(29)	69.54±0.73(28)	75.78±1.08(24)	64.5±4.62(32)	83.19±0.66(12)	83.26±0.00(1)	84.5±0.00(1)	80.74±1.03(4)	77.29±3.52(20)
mammography	41.76±0.02(6)	41.65±0.00(7)	18.52±9.52(29)	21.32±0.00(26)	37.08±2.83(9.5)	27.82±3.84(22)	11.77±0.01(31)	54.63 ±0.00(2)	55.2 ±0.00(1)	27.54±1.44(23)	22.01±1.50(44)	22.01±1.50(44)	27.24±2.23(24)	41.04±0.00(1)	39.83±0.00(1)	28.51±2.57(21)	28.91±1.92(20)
mnist	64.99±0.00(12.5)	64.99±0.00(12.5)	55.22±3.33(21)	22.21±0.00(29)	47.97±4.73(23)	65.09±0.57(11)	20.19±0.00(30)	16.86±0.03(1.5)	16.86±0.03(1.5)	46.03±5.58(25)	34.07±7.67(27)	46.06±7.88(24)	59.72±1.98(15)	72.08±0.00(3)	70.36±0.00(5)	56.56±6.65(17)	41.15±1.11(26)
mnist-coverage	100.0±0.00(8.5)	100.0±0.00(8.5)	32.68±3.38(31)	100.0 ±0.00(8.5)	100.0 ±0.00(8.5)	100.0 ±0.00(8.5)	72.21±0.00(30)	98.2±0.00(18)	98.2±0.00(18)	99.91±0.17(17)	90.8±10.84(23)	70.61±23.94(27)	15.65±19.61(13)	100.0 ±0.00(8.5)	100.0 ±0.00(8.5)	89.48±1.94(24)	62.46±0.00(29)
music	6.02±0.00(24)	6.02±0.00(24)	3.93±0.47(31.5)	11.57±1.15(21)	11.57±1.15(21)	5.14±0.00(28)	5.14±0.00(28)	5.59±0.26(35)	5.59±0.26(35)	4.53±0.40(30)	3.93±0.47(31.5)	4.95±2.39(29)	21.18±5.39(41)	27.11±0.00(1)	21.18±5.39(41)	41.08±1.33(45)	10.77±1.02(20)
opendigits	59.39±0.00(20)	59.39±0.00(20)	58.26±0.00(21)	42.38±0.00(3)	46.08±0.67(29)	63.51±1.71(3)	59.1±0.00(2)	41.51±0.00(31)	38.54±0.00(23)	42.05±3.89(27)	48.57±3.88(28)	60.26±12.84(18)	73.46 ±2.81(1)	61.76±0.00(1)	67.15±0.00(8)	64.91±1.13(10)	54.53±1.01(26)
paghebooks	39.14±0.00(19)	38.63±0.00(20)	14.47±4.88(29)	42.33±0.00(4)	14.62±1.53(927)	33.35±2.85(23)	22.26±0.00(26)	30.86±0.00(24)	41.45±0.00(18)	9.34±7.78(32)	37.23±7.94(21)	11.71±9.83(31)	14.57±3.34(32)	95.99 ±0.00(2)	82.34±0.00(6)	58.33±6.17(13)	29.79±0.96(25)
pedigree	71.18±3.59(15.5)	71.18±3.59(15.5)	17.22±2.96(14)	82.88±2.36(6)	61.64±6.66(27)	65.75±8.84(2)	86.78±1.69(30)	60.97±2.47(19)	64.77±2.42(25)	93.75±1.72(28)	59.36±7.09(29)	56.48±5.53(31)	53.42±1.53(52)	75.31±2.22(5)	74.01±2.75(9)	75.31±2.22(5)	67.61±2.81(25)
pin	82.94±0.28(13)	82.94±0.28(13)	62.47±5.19(28)	87.68±0.00(6)	80.25±1.95(23)	95.89±0.18(2)	82.94±0.28(13)	85.72±0.00(1)	95.89±0.18(2)	76.38±2.12(9)	63.72±0.69(12)	47.84±1.34(30)	79.32±1.88(25)	96.02±0.00(3)	97.32 ±0.00(7)	71.56±1.98(27)	49.34±5.54(31)
satellite	96.27±0.00(9.5)	96.27±0.00(9.5)	91.66±1.93(28)	97.49±0.00(16)	99.84±4.72(4)	60.18±2.92(28)	94.87±0.00(2)	98.05±0.00(1)	95.24±0.00(2)	98.03±0.13(2)	55.74±0.66(29)	65.98±23.68(27)	13.35±1.88(23)	96.37±0.00(7)	92.71±0.16(5)	96.37±0.00(7)	76.78±3.46(26)
shirt	40.14±0.33(27)	36.39±0.33(28)	74.74±17.48(3)	53.37±0.59(20)	31.88±2.23(20)	42.18±1.84(26)	63.01±1.72(15)	29.09±0.18(32)	30.49±0.23(1)	42.99±3.28(25)	53.04±5.71(21)	50.37±1.79(23)	65.62±1.18(13)	96.37±0.00(7)	92.71±0.16(5)	56.71±0.16(5)	62.54±0.34(20)
skin	49.38±6.44(13)	49.5±6.11(1)	0.77±0.40(30)	1.15±0.12(7)	0.1±0.00(32)	32.48±4.81(29)	49.46±6.64(12)	0.99±0.00(29)	0.99±0.00(29)	30.73±22.82(21)	8.16±5.47(24)	20.92±2.56(32)	8.69±19.27(23)	53.94±7.83(4)	50.27±5.76(8)	38.67±3.38(16)	31.76±4.31(20)
snub	2.77±0.00(28.5)	2.77±0.00(28.5)	81.84±0.00(15.5)	85.36±2.63(4)	78.42±0.00(23)	82.09±0.21(13)	50.5±0.00(2)	73.8±0.00(29)	71.26±0.00(2)	75.26±2.44(24)	80.16±4.45(20)	74.22±2.55(25)	79.07±3.12(1)	83.01±0.00(1)	82.54±0.00(1)	73.08±0.13(27)	72.9±0.52(20)
speech	2.77±0.00(28.5)	2.77±0.00(28.5)	81.84±0.00(15.5)	85.36±2.63(4)	78.42±0.00(23)	82.09±0.21(13)	50.5±0.00(2)	73.8±0.00(29)	71.26±0.00(2)	75.26±2.44(24)	80.16±4.45(20)	74.22±2.55(25)	79.07±3.12(1)	83.01±0.00(1)	82.54±0.00(1)	73.08±0.13(27)	72.9±0.52(20)
stamps	59.92±7.99(15)	58.81±7.82(17)	52.41±12.56(21)	52.28±4.56(22)	33.47±0.36(31)	49.13±8.19(26)	56.43±1.20(4)	49.03±8.62(7)	42.62±9.94(29)	57.17±1.21(19)	54.72±2.11(28)	28.48±21.94(32)	72.02±4.82(6)	72.02±4.82(6)	65.57±1.10(6)	50.81±7.62(23)	3.09±0.15(13)
thyroid	81.33±0.06(0)	81.34±0.00(5)	75.79±6.78(15)	76.95±0.00(13)	53.61±2.83(128)	80.09±0.89(9)	60.91±0.99(2)	30.19±0.00(3)	40.03±0.00(2)	49.06±8.11(8)	64.26±6.24(21)	63.08±15.77(23)	23.35±10.41(9)	81.47±0.00(4)	81.47±0.00(4)	66.61±0.60(20)	66.61±0.60(20

Table 18.1: 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.

Dataset	FoMe-OD (d=100)	FoMe-OD (d=20)	DTE-SP	INN	KL	DTE-C	LoF	CBLoF	FeatureBagging	SLAD	DPIM	OC-SVM	DTE-IG	IRForest	MCD
air	7.82 $\pm$ 0.41(4)	6.49 $\pm$ 0.71(12)	5.82 $\pm$ 0.07(17)	5.9 $\pm$ 0.01(5)	4.91 $\pm$ 0.56(22)	4.2 $\pm$ 0.26(5)	8.16 $\pm$ 0.00(3)	6.74 $\pm$ 0.08(10)	5.93 $\pm$ 0.57(2)	5.32 $\pm$ 0.01(19)	6.76 $\pm$ 0.19(9)	7.29 $\pm$ 0.00(8)	5.12 $\pm$ 0.06(821)	4.2 $\pm$ 0.26(26.5)	3.41 $\pm$ 0.14(31)
amazon	11.72 $\pm$ 3.24(4)	10.72 $\pm$ 1.21(6)	10.8 $\pm$ 0.01(5)	11.4 $\pm$ 0.07(5)	9.4 $\pm$ 0.38(28)	11.78 $\pm$ 1.53(2)	10.52 $\pm$ 0.02(3)	10.52 $\pm$ 0.02(3)	10.05 $\pm$ 0.74(2)	10.16 $\pm$ 0.02(21)	11.08 $\pm$ 0.11(11)	12.0 $\pm$ 0.00(1)	10.88 $\pm$ 2.37(20)	11.28 $\pm$ 0.64(10)	11.32 $\pm$ 0.23(9)
anthyroid	49.25 $\pm$ 0.00(2)	51.35 $\pm$ 0.01(7)	61.84 $\pm$ 1.59(4)	61.99 $\pm$ 0.00(5)	49.44 $\pm$ 3.88(28)	47.78 $\pm$ 1.50(1)	46.6 $\pm$ 0.00(2)	56.73 $\pm$ 3.31(2)	56.73 $\pm$ 3.31(2)	56.73 $\pm$ 3.31(2)	57.23 $\pm$ 2.82(10)	57.23 $\pm$ 2.82(10)	57.23 $\pm$ 2.82(10)	55.02 $\pm$ 4.22(14)	50.7 $\pm$ 0.01(9)
backdoor	96.46 $\pm$ 0.72(1)	96.46 $\pm$ 0.72(1)	96.46 $\pm$ 0.72(1)	96.46 $\pm$ 0.72(1)	96.46 $\pm$ 0.72(1)	96.46 $\pm$ 0.72(1)	96.46 $\pm$ 0.72(1)	96.46 $\pm$ 0.72(1)	96.46 $\pm$ 0.72(1)	96.46 $\pm$ 0.72(1)	96.46 $\pm$ 0.72(1)	96.46 $\pm$ 0.72(1)	96.46 $\pm$ 0.72(1)	96.46 $\pm$ 0.72(1)	96.46 $\pm$ 0.72(1)
breast	95.92 $\pm$ 0.72(1)	95.92 $\pm$ 0.72(1)	95.92 $\pm$ 0.72(1)	95.92 $\pm$ 0.72(1)	95.92 $\pm$ 0.72(1)	95.92 $\pm$ 0.72(1)	95.92 $\pm$ 0.72(1)	95.92 $\pm$ 0.72(1)	95.92 $\pm$ 0.72(1)	95.92 $\pm$ 0.72(1)	95.92 $\pm$ 0.72(1)	95.92 $\pm$ 0.72(1)	95.92 $\pm$ 0.72(1)	95.92 $\pm$ 0.72(1)	95.92 $\pm$ 0.72(1)
campain	90.61 $\pm$ 0.61(4)	90.61 $\pm$ 0.61(4)	90.61 $\pm$ 0.61(4)	90.61 $\pm$ 0.61(4)	90.61 $\pm$ 0.61(4)	90.61 $\pm$ 0.61(4)	90.61 $\pm$ 0.61(4)	90.61 $\pm$ 0.61(4)	90.61 $\pm$ 0.61(4)	90.61 $\pm$ 0.61(4)	90.61 $\pm$ 0.61(4)	90.61 $\pm$ 0.61(4)	90.61 $\pm$ 0.61(4)	90.61 $\pm$ 0.61(4)	90.61 $\pm$ 0.61(4)
cardio	72.16 $\pm$ 0.05(5)	68.86 $\pm$ 2.67(9)	67.7 $\pm$ 0.01(3)	61.93 $\pm$ 0.01(7)	52.16 $\pm$ 2.89(25)	53.37 $\pm$ 2.62(23)	62.5 $\pm$ 0.00(5)	70.93 $\pm$ 0.04(8)	70.93 $\pm$ 0.04(8)	70.93 $\pm$ 0.04(8)	70.93 $\pm$ 0.04(8)	70.93 $\pm$ 0.04(8)	70.93 $\pm$ 0.04(8)	70.93 $\pm$ 0.04(8)	70.93 $\pm$ 0.04(8)
cardiography	56.22 $\pm$ 0.05(4)	48.41 $\pm$ 1.52(13.5)	46.78 $\pm$ 1.44(19)	46.35 $\pm$ 0.02(9)	38.33 $\pm$ 2.89(25)	38.33 $\pm$ 2.89(25)	42.44 $\pm$ 0.02(0)	42.44 $\pm$ 0.02(0)	42.44 $\pm$ 0.02(0)	42.44 $\pm$ 0.02(0)	42.44 $\pm$ 0.02(0)	42.44 $\pm$ 0.02(0)	42.44 $\pm$ 0.02(0)	42.44 $\pm$ 0.02(0)	42.44 $\pm$ 0.02(0)
cellula	8.41 $\pm$ 0.88(28)	7.38 $\pm$ 1.98(29)	15.81 $\pm$ 0.69(18)	17.19 $\pm$ 0.83(16)	12.69 $\pm$ 1.83(22)	17.35 $\pm$ 3.48(14)	19.1 $\pm$ 0.47(2)	25.32 $\pm$ 7.14(5)	29.23 $\pm$ 0.93(1)	31.69 $\pm$ 1.66(20)	26.02 $\pm$ 2.89(4)	27.97 $\pm$ 0.74(1)	19.11 $\pm$ 5.61(10.5)	56.14 $\pm$ 2.63(1)	17.33 $\pm$ 2.29(15)
ceus	69.92 $\pm$ 7.24(1)	66.84 $\pm$ 1.32(1)	66.84 $\pm$ 1.32(1)	66.84 $\pm$ 1.32(1)	66.84 $\pm$ 1.32(1)	66.84 $\pm$ 1.32(1)	66.84 $\pm$ 1.32(1)	66.84 $\pm$ 1.32(1)	66.84 $\pm$ 1.32(1)	66.84 $\pm$ 1.32(1)	66.84 $\pm$ 1.32(1)	66.84 $\pm$ 1.32(1)	66.84 $\pm$ 1.32(1)	66.84 $\pm$ 1.32(1)	66.84 $\pm$ 1.32(1)
comp	69.92 $\pm$ 7.24(1)	66.84 $\pm$ 1.32(1)	66.84 $\pm$ 1.32(1)	66.84 $\pm$ 1.32(1)	66.84 $\pm$ 1.32(1)	66.84 $\pm$ 1.32(1)	66.84 $\pm$ 1.32(1)	66.84 $\pm$ 1.32(1)	66.84 $\pm$ 1.32(1)	66.84 $\pm$ 1.32(1)	66.84 $\pm$ 1.32(1)	66.84 $\pm$ 1.32(1)	66.84 $\pm$ 1.32(1)	66.84 $\pm$ 1.32(1)	66.84 $\pm$ 1.32(1)
donors	97.85 $\pm$ 0.79(1)	77.33 $\pm$ 3.22(10)	62.9 $\pm$ 2.83(6)	94.91 $\pm$ 0.67(3)	97.22 $\pm$ 1.04(2)	82.17 $\pm$ 2.48(8)	74.47 $\pm$ 2.04(1)	48.48 $\pm$ 1.33(14)	48.48 $\pm$ 1.33(14)	48.48 $\pm$ 1.33(14)	48.48 $\pm$ 1.33(14)	48.48 $\pm$ 1.33(14)	48.48 $\pm$ 1.33(14)	48.48 $\pm$ 1.33(14)	48.48 $\pm$ 1.33(14)
fault	57.35 $\pm$ 0.08(1)	55.39 $\pm$ 3.22(20)	56.2 $\pm$ 0.41(6)	55.57 $\pm$ 0.01(9)	57.99 $\pm$ 0.56(7)	56.23 $\pm$ 1.73(15)	50.67 $\pm$ 0.01(3)	56.4 $\pm$ 0.87(12)	67.65 $\pm$ 4.19(3)	47.39 $\pm$ 4.57(15)	58.45 $\pm$ 1.15(5)	35.13 $\pm$ 0.52(21)	55.11 $\pm$ 0.96(12)	53.64 $\pm$ 1.34(25.5)	33.32 $\pm$ 1.42(125)
frad	61.34 $\pm$ 11.86(6)	44.63 $\pm$ 8.94(20)	48.36 $\pm$ 5.41(13)	45.22 $\pm$ 4.87(18)	57.83 $\pm$ 5.97(11)	67.23 $\pm$ 13.11(2)	39.74 $\pm$ 3.16(17)	23.79 $\pm$ 0.61(26)	22.45 $\pm$ 7.84(30)	35.09 $\pm$ 4.75(15)	72.16 $\pm$ 1.29(4)	14.97 $\pm$ 7.84(30)	78.1 $\pm$ 13.6(2)	16.18 $\pm$ 7.01(26)	56.19 $\pm$ 3.88(11)
glass	24.88 $\pm$ 16.83(15)	59.96 $\pm$ 10.09(6)	24.88 $\pm$ 16.83(15)	24.88 $\pm$ 16.83(15)	24.88 $\pm$ 16.83(15)	24.88 $\pm$ 16.83(15)	24.88 $\pm$ 16.83(15)	24.88 $\pm$ 16.83(15)	24.88 $\pm$ 16.83(15)	24.88 $\pm$ 16.83(15)	24.88 $\pm$ 16.83(15)	24.88 $\pm$ 16.83(15)	24.88 $\pm$ 16.83(15)	24.88 $\pm$ 16.83(15)	24.88 $\pm$ 16.83(15)
hepatitis	99.84 $\pm$ 0.71(2.5)	97.4 $\pm$ 3.76(1)	79.03 $\pm$ 11.95(14)	81.29 $\pm$ 5.04(11)	99.64 $\pm$ 0.70(29)	92.8 $\pm$ 3.32(8)	61.9 $\pm$ 0.62(29)	66.91 $\pm$ 8.53(16)	42.57 $\pm$ 1.46(6)	49.94 $\pm$ 0.79(2.5)	86.69 $\pm$ 4.26(9)	99.84 $\pm$ 0.71(2.5)	99.84 $\pm$ 0.71(2.5)	99.84 $\pm$ 0.71(2.5)	99.84 $\pm$ 0.71(2.5)
hepatit	99.84 $\pm$ 0.71(2.5)	97.4 $\pm$ 3.76(1)	79.03 $\pm$ 11.95(14)	81.29 $\pm$ 5.04(11)	99.64 $\pm$ 0.70(29)	92.8 $\pm$ 3.32(8)	61.9 $\pm$ 0.62(29)	66.91 $\pm$ 8.53(16)	42.57 $\pm$ 1.46(6)	49.94 $\pm$ 0.79(2.5)	86.69 $\pm$ 4.26(9)	99.84 $\pm$ 0.71(2.5)	99.84 $\pm$ 0.71(2.5)	99.84 $\pm$ 0.71(2.5)	99.84 $\pm$ 0.71(2.5)
ind	10.12 $\pm$ 4.26(6)	7.61 $\pm$ 1.81(0)	5.2 $\pm$ 0.02(8)	5.44 $\pm$ 0.00(26)	10.56 $\pm$ 0.94(2)	7.24 $\pm$ 1.62(12)	64.4 $\pm$ 0.01(8)	6.96 $\pm$ 0.90(45)	6.56 $\pm$ 0.07(1)	10.24 $\pm$ 0.42(1)	5.56 $\pm$ 0.09(25)	5.84 $\pm$ 0.02(1)	10.66 $\pm$ 1.53(3)	6.2 $\pm$ 0.49(20)	7.44 $\pm$ 0.21(1)
interests	40.82 $\pm$ 1.48(25)	33.83 $\pm$ 7.72(7)	33.83 $\pm$ 7.72(7)	51.9 $\pm$ 0.01(3)	55.67 $\pm$ 0.91(4)	67.8 $\pm$ 2.48(1)	54.6 $\pm$ 0.00(8)	45.76 $\pm$ 0.24(1)	54.35 $\pm$ 0.33(2)	57.63 $\pm$ 0.32(1)	88.64 $\pm$ 1.51(5)	92.62 $\pm$ 1.51(5)	89.88 $\pm$ 1.21(3)	83.43 $\pm$ 3.52(2)	88.64 $\pm$ 1.21(5.5)
ionosphere	92.31 $\pm$ 1.33(6)	89.99 $\pm$ 1.04(12)	91.51 $\pm$ 2.09(8)	90.47 $\pm$ 1.17(11)	94.18 $\pm$ 1.62(1)	94.18 $\pm$ 1.62(1)	94.18 $\pm$ 1.62(1)	91.96 $\pm$ 2.14(7)	87.65 $\pm$ 2.26(18)	92.61 $\pm$ 1.26(4)	88.64 $\pm$ 1.51(5)	92.62 $\pm$ 1.51(5)	89.88 $\pm$ 1.21(3)	83.43 $\pm$ 3.52(2)	88.64 $\pm$ 1.21(5.5)
landsat	52.66 $\pm$ 0.04(4)	49.93 $\pm$ 1.75(9)	51.22 $\pm$ 2.57(7)	51.46 $\pm$ 0.00(6)	53.82 $\pm$ 0.35(2)	34.29 $\pm$ 2.97(23)	33.54 $\pm$ 0.03(3)	38.33 $\pm$ 0.19(22)	47.33 $\pm$ 4.45(3)	46.03 $\pm$ 0.04(12)	40.23 $\pm$ 1.02(19)	50.69 $\pm$ 0.02(18)	30.26 $\pm$ 4.37(32)	43.27 $\pm$ 1.31(4)	47.73 $\pm$ 9.58(11)
letter	4.0 $\pm$ 0.01(5)	0.23 $\pm$ 0.43(2)	1.0 $\pm$ 0.02(7)	1.03 $\pm$ 0.02(7)	1.03 $\pm$ 0.02(7)	1.03 $\pm$ 0.02(7)	1.03 $\pm$ 0.02(7)	1.03 $\pm$ 0.02(7)	1.03 $\pm$ 0.02(7)	1.03 $\pm$ 0.02(7)	1.03 $\pm$ 0.02(7)	1.03 $\pm$ 0.02(7)	1.03 $\pm$ 0.02(7)	1.03 $\pm$ 0.02(7)	1.03 $\pm$ 0.02(7)
lymphography	96.39 $\pm$ 5.14(7)	95.78 $\pm$ 5.81(8)	95.78 $\pm$ 5.81(8)	95.78 $\pm$ 5.81(8)	95.78 $\pm$ 5.81(8)	95.78 $\pm$ 5.81(8)	95.78 $\pm$ 5.81(8)	95.78 $\pm$ 5.81(8)	95.78 $\pm$ 5.81(8)	95.78 $\pm$ 5.81(8)	95.78 $\pm$ 5.81(8)	95.78 $\pm$ 5.81(8)	95.78 $\pm$ 5.81(8)	95.78 $\pm$ 5.81(8)	95.78 $\pm$ 5.81(8)
lympho	96.39 $\pm$ 5.14(7)	95.78 $\pm$ 5.81(8)	95.78 $\pm$ 5.81(8)	95.78 $\pm$ 5.81(8)	95.78 $\pm$ 5.81(8)	95.78 $\pm$ 5.81(8)	95.78 $\pm$ 5.81(8)	95.78 $\pm$ 5.81(8)	95.78 $\pm$ 5.81(8)	95.78 $\pm$ 5.81(8)	95.78 $\pm$ 5.81(8)	95.78 $\pm$ 5.81(8)	95.78 $\pm$ 5.81(8)	95.78 $\pm$ 5.81(8)	95.78 $\pm$ 5.81(8)
manometry	37.04 $\pm$ 0.62(3)	42.38 $\pm$ 1.00(12)	42.38 $\pm$ 1.00(12)	40.38 $\pm$ 0.01(2)	17.88 $\pm$ 3.62(29)	3.71 $\pm$ 3.91(6)	38.46 $\pm$ 0.01(3)	39.38 $\pm$ 0.01(3)	39.38 $\pm$ 0.01(3)	22.15 $\pm$ 1.26(27)	24.62 $\pm$ 0.41(26)	41.92 $\pm$ 0.01(1)	35.11 $\pm$ 7.06(20)	30.23 $\pm$ 2.48(14)	2.62 $\pm$ 0.03(2)
mnist	62.04 $\pm$ 0.01(4)	31.71 $\pm$ 4.25(28)	72.6 $\pm$ 1.21(1)	71.86 $\pm$ 0.00(2)	54.46 $\pm$ 2.88(16)	71.43 $\pm$ 0.00(3)	100.0 $\pm$ 0.00(8.5)	100.0 $\pm$ 0.00(8.5)	68.89 $\pm$ 1.44(9)	67.61 $\pm$ 0.39(7)	67.37 $\pm$ 3.89(15)	100.0 $\pm$ 0.00(8.5)	50.43 $\pm$ 8.41(22)	52.6 $\pm$ 1.61(18)	55.97 $\pm$ 2.61(18)
mnist	93.23 $\pm$ 3.23(18)	92.16 $\pm$ 9.17(20)	100.0 $\pm$ 0.00(8.5)	100.0 $\pm$ 0.00(8.5)	83.34 $\pm$ 8.97(25)	100.0 $\pm$ 0.00(8.5)	100.0 $\pm$ 0.00(8.5)	100.0 $\pm$ 0.00(8.5)	68.89 $\pm$ 1.44(9)	67.61 $\pm$ 0.39(7)	67.37 $\pm$ 3.89(15)	100.0 $\pm$ 0.00(8.5)	88.66 $\pm$ 25.36(23)	35.88 $\pm$ 24.78(30)	55.97 $\pm$ 2.61(18)
optdigits	37.33 $\pm$ 0.07(1)	27.23 $\pm$ 0.73(9)	27.23 $\pm$ 0.73(9)	21.93 $\pm$ 0.01(2)	57.73 $\pm$ 8.41(1)	109.33 $\pm$ 3.39(16)	109.33 $\pm$ 3.39(16)	109.33 $\pm$ 3.39(16)	47.33 $\pm$ 4.45(3)	39.87 $\pm$ 0.07(8)	28.43 $\pm$ 7.08(8)	0.67 $\pm$ 0.02(15)	27.07 $\pm$ 18.57(10)	12.8 $\pm$ 6.28(15)	53.91 $\pm$ 14.3(29)
pageblocks	61.18 $\pm$ 0.00(10)	64.71 $\pm$ 0.00(5)	59.29 $\pm$ 2.81(2)	59.02 $\pm$ 0.01(3)	64.91 $\pm$ 1.29(4)	62.12 $\pm$ 0.47(8)	63.88 $\pm$ 0.00(2)	65.29 $\pm$ 0.28(3)	63.45 $\pm$ 1.67(6)	60.21 $\pm$ 0.01(1)	50.31 $\pm$ 0.71(22)	55.09 $\pm$ 0.01(8)	54.99 $\pm$ 2.54(20)	42.63 $\pm$ 2.29(26)	57.57 $\pm$ 0.11(7)
pendigit	70.55 $\pm$ 0.00(10)	70.55 $\pm$ 0.00(10)	70.55 $\pm$ 0.00(10)	70.55 $\pm$ 0.00(10)	70.55 $\pm$ 0.00(10)	70.55 $\pm$ 0.00(10)	70.55 $\pm$ 0.00(10)	70.55 $\pm$ 0.00(10)	70.55 $\pm$ 0.00(10)	70.55 $\pm$ 0.00(10)	70.55 $\pm$ 0.00(10)	70.55 $\pm$ 0.00(10)	70.55 $\pm$ 0.00(10)	70.55 $\pm$ 0.00(10)	70.55 $\pm$ 0.00(10)
pendigit	70.55 $\pm$ 0.00(10)	70.55 $\pm$ 0.00(10)	70.55 $\pm$ 0.00(10)	70.55 $\pm$ 0.00(10)	70.55 $\pm$ 0.00(10)	70.55 $\pm$ 0.00(10)	70.55 $\pm$ 0.00(10)	70.55 $\pm$ 0.00(10)	70.55 $\pm$ 0.00(10)	70.55 $\pm$ 0.00(10)	70.55 $\pm$ 0.00(10)	70.55 $\pm$ 0.00(10)	70.55 $\pm$ 0.00(10)	70.55 $\pm$ 0.00(10)	70.55 $\pm$ 0.00(10)
satellite	77.25 $\pm$ 0.00(3)	74.26 $\pm$ 0.66(9)	71.91 $\pm$ 0.98(12)	71.81 $\pm$ 0.00(13)	74.99 $\pm$ 0.46(5)	72.33 $\pm$ 0.63(1)	72.33 $\pm$ 0.63(1)	64.02 $\pm$ 0.07(26)	72.64 $\pm$ 0.08(10)	78.24 $\pm$ 0.13(2)	73.74 $\pm$ 0.27(1)	67.34 $\pm$ 0.00(20)	65.71 $\pm$ 5.09(15)	67.12 $\pm$ 0.64(21)	63.15 $\pm$ 1.11(29)
satimgae-2	88.73 $\pm$ 0.01(11.5)	88.73 $\pm$ 0.01(11.5)	90.14 $\pm$ 0.07(5)	90.14 $\pm$ 0.07(5)	88.45 $\pm$ 1.50(14)	66.48 $\pm$ 2.71(28)	81.69 $\pm$ 0.01(9)	92.96 $\pm$ 0.01(2)	84.51 $\pm$ 1.01(7)	88.73 $\pm$ 0.01(11.5)	78.59 $\pm$ 5.12(22)	91.55 $\pm$ 0.04(4)	78.31 $\pm$ 5.23(23)	96.71 $\pm$ 0.51(6)	95.77 $\pm$ 0.01(1)
shuttle	97.72 $\pm$ 0.21(5)	98.3 $\$													

Table 18.2: 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.

Dataset	VAE	PCA	PlanFlow	HBOs	GANomaly	GOAD	DIF	COPOD	ECOD	DeepSVDD	LODA	DGMM	DRDCC	DTE- NP^{*}	KN N^{*}	QL *	DTE *
amazon	7.63±0.005 (5)	7.63±0.005 (5)	3.83±0.72 (29)	9.35±0.007 (7)	9.35±2.27 (1)	5.73±1.45 (18)	9.31±0.028 (8)	4.33±0.024 (8)	4.44±0.025 (8)	71.7±0.02 (20)	10.6±1.58 (11)	5.98±1.67 (14)	0.0±0.0 (32)	5.86±0.00 (8)	6.01±0.00 (13)	4.99±0.07 (23)	3.56±0.03 (30)
amazon	11.03±0.002 (5)	11.03±0.002 (5)	9.86±0.74 (30)	10.63±0.00 (18)	8.96±2.47 (10)	5.15±0.26 (5)	9.32±0.06 (29)	4.38±0.024 (8)	4.40±0.023 (8)	11.76±2.19 (3)	60.6±1.18 (17)	6.04±1.18 (17)	0.0±0.0 (32)	10.6±0.00 (16)	10.92±0.00 (15)	10.92±0.00 (15)	10.92±0.00 (15)
amazon	5.01±0.002 (1)	5.01±0.002 (1)	6.63±0.77 (27)	6.93±0.06 (25)	2.19±0.47 (46)	4.79±2.47 (32)	20.28±2.02 (17)	0.0±0.0 (31)	0.0±0.0 (31)	82.96±3.28 (5)	46.94±4.18 (28)	5.25±3.94 (26)	57.42±2.53 (10)	42.86±1.32 (12)	31.87±1.12 (15)	86.76±1.0 (2)	47.75±1.54 (13)
breast	9.18±0.004 (5)	9.18±0.004 (5)	9.42±0.16 (62)	9.63±0.34 (42)	9.05±3.37 (25)	22.62±0.05 (1)	5.81±1.45 (83)	96.41±0.34 (8)	96.43±0.35 (21)	37.89±1.29 (25)	30.74±5.47 (18)	34.13±3.43 (27)	48.27±2.58 (32)	50.23±0.00 (5)	95.97±0.31 (11)	50.07±0.49 (8)	92.35±0.79 (25)
breast	48.41±0.04 (4)	48.41±0.04 (4)	49.71±0.04 (7)	40.92±4.17 (1)	40.92±4.17 (1)	22.62±0.05 (1)	5.81±1.45 (83)	96.41±0.34 (8)	96.43±0.35 (21)	37.89±1.29 (25)	30.74±5.47 (18)	34.13±3.43 (27)	48.27±2.58 (32)	50.23±0.00 (5)	95.97±0.31 (11)	50.07±0.49 (8)	92.35±0.79 (25)
cardio	76.18±0.04 (5)	76.18±0.04 (5)	59.77±1.54 (20)	56.25±0.00 (26)	46.75±3.47 (22)	59.89±0.03 (5)	27.27±0.03 (2)	27.27±0.03 (2)	27.27±0.03 (2)	38.41±4.19 (29)	38.41±4.19 (29)	38.41±4.19 (29)	46.93±2.34 (28)	46.93±2.34 (28)	46.93±2.34 (28)	34.91±0.06 (31)	34.91±0.06 (31)
cardiography	91.59±0.002 (5)	91.59±0.002 (5)	94.74±0.33 (11)	71.42±0.00 (2)	71.42±0.00 (2)	59.89±0.03 (5)	27.27±0.03 (2)	27.27±0.03 (2)	27.27±0.03 (2)	38.41±4.19 (29)	38.41±4.19 (29)	38.41±4.19 (29)	46.93±2.34 (28)	46.93±2.34 (28)	46.93±2.34 (28)	34.91±0.06 (31)	34.91±0.06 (31)
cellula	7.04±0.001 (5)	7.04±0.001 (5)	7.19±1.49 (12)	7.19±1.49 (12)	7.19±1.49 (12)	7.19±1.49 (12)	7.19±1.49 (12)	7.19±1.49 (12)	7.19±1.49 (12)	11.81±0.00 (1)	11.81±0.00 (1)	11.81±0.00 (1)	11.81±0.00 (1)	11.81±0.00 (1)	11.81±0.00 (1)	11.81±0.00 (1)	11.81±0.00 (1)
cellula	7.04±0.001 (5)	7.04±0.001 (5)	7.19±1.49 (12)	7.19±1.49 (12)	7.19±1.49 (12)	7.19±1.49 (12)	7.19±1.49 (12)	7.19±1.49 (12)	7.19±1.49 (12)	11.81±0.00 (1)	11.81±0.00 (1)	11.81±0.00 (1)	11.81±0.00 (1)	11.81±0.00 (1)	11.81±0.00 (1)	11.81±0.00 (1)	11.81±0.00 (1)
cos	16.24±1.68 (22)	16.24±1.68 (22)	22.91±2.48 (30)	10.75±1.26 (20)	10.75±1.26 (20)	4.00±0.0 (32)	1.19±0.88 (31)	18.82±0.00 (1)	24.66±1.11 (8)	3.43±3.44 (29)	24.91±1.26 (19)	12.66±1.56 (24)	41.87±5.75 (12)	61.65±1.92 (11)	52.19±1.92 (11)	37.61±2.80 (15)	38.79±0.02 (14)
cos	16.24±1.68 (22)	16.24±1.68 (22)	22.91±2.48 (30)	10.75±1.26 (20)	10.75±1.26 (20)	4.00±0.0 (32)	1.19±0.88 (31)	18.82±0.00 (1)	24.66±1.11 (8)	3.43±3.44 (29)	24.91±1.26 (19)	12.66±1.56 (24)	41.87±5.75 (12)	61.65±1.92 (11)	52.19±1.92 (11)	37.61±2.80 (15)	38.79±0.02 (14)
donors	37.75±0.93 (22)	37.75±0.93 (22)	47.64±1.58 (15)	35.64±0.68 (28)	18.98±1.64 (73)	4.29±1.43 (32)	6.41±0.83 (26)	91.82±0.00 (2)	91.82±0.00 (2)	44.92±0.41 (9)	21.04±2.71 (54)	21.04±2.71 (54)	56.73±2.54 (26)	56.32±0.00 (4)	86.63±0.08 (7)	78.24±2.40 (16)	72.80±0.47 (16)
donors	37.75±0.93 (22)	37.75±0.93 (22)	47.64±1.58 (15)	35.64±0.68 (28)	18.98±1.64 (73)	4.29±1.43 (32)	6.41±0.83 (26)	91.82±0.00 (2)	91.82±0.00 (2)	44.92±0.41 (9)	21.04±2.71 (54)	21.04±2.71 (54)	56.73±2.54 (26)	56.32±0.00 (4)	86.63±0.08 (7)	78.24±2.40 (16)	72.80±0.47 (16)
fland	34.06±3.22 (25)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)
fland	34.06±3.22 (25)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)
fland	34.06±3.22 (25)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)
fland	34.06±3.22 (25)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)
fland	34.06±3.22 (25)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)
fland	34.06±3.22 (25)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)
fland	34.06±3.22 (25)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)
fland	34.06±3.22 (25)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)
fland	34.06±3.22 (25)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)
fland	34.06±3.22 (25)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)
fland	34.06±3.22 (25)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)
fland	34.06±3.22 (25)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)
fland	34.06±3.22 (25)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)
fland	34.06±3.22 (25)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)
fland	34.06±3.22 (25)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)
fland	34.06±3.22 (25)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)
fland	34.06±3.22 (25)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)
fland	34.06±3.22 (25)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)
fland	34.06±3.22 (25)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)
fland	34.06±3.22 (25)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)
fland	34.06±3.22 (25)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)
fland	34.06±3.22 (25)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)
fland	34.06±3.22 (25)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)
fland	34.06±3.22 (25)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)
fland	34.06±3.22 (25)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)
fland	34.06±3.22 (25)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)	33.76±1.67 (27)
fland																	

N. Benchmark OD Datasets

Table 19: 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

O. 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.

P. 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 semi-supervised OD with clean/inlier-only training data. Future work can study the unsupervised OD setting and pretraining with mixed/“contaminated” data in this 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 the first foundation model for OD, FoMo-OD inspires many promising directions for future research that could lead to fruition for additional practical applications.

Q. 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://anonymous.4open.science/r/PFN40D>. Further, full implementation details are provided in Appendix F.