

Unsupervised Surrogate Anomaly Detection

Simon Klüttermann^{1,3}[0000-0001-9698-4339], Tim Katzke^{1,2}[0009-0000-0154-7735],
and Emmanuel Müller^{1,2,3}[0000-0002-5409-6875]

¹ TU Dortmund University, Germany

² Research Center Trustworthy Data Science and Security

³ Lamarr Institute for Machine Learning and Artificial Intelligence
`firstname.lastname@cs.tu-dortmund.de`

Abstract. In this paper, we study unsupervised anomaly detection algorithms that learn a neural network representation, i.e. regular patterns of normal data, which anomalies are deviating from. Inspired by a similar concept in engineering, we refer to our methodology as surrogate anomaly detection. We formalize the concept of surrogate anomaly detection into a set of axioms required for optimal surrogate models and propose a new algorithm, named DEAN (Deep Ensemble ANomaly detection), designed to fulfill these criteria. We evaluate DEAN on 121 benchmark datasets, demonstrating its competitive performance against 19 existing methods, as well as the scalability and reliability of our method.

Keywords: Anomaly Detection · Ensemble Methods.

1 Introduction

Anomaly detection (AD) is a crucial subdomain of data analytics with countless applications ranging from fraud detection [24] to health monitoring [10]. In all of these domains, anomalies are rare, exceptional or interesting objects that are highly deviating from the residual (regular) data [46]. Generally, anomaly detection can be approached in three primary ways: with extensive access to labeled anomalies (supervised), with access to a limited number of labeled anomalies (semi-supervised), or without labeled anomalies (unsupervised). In this paper, we focus on unsupervised anomaly detection. This setting poses the unique challenge of identifying a wide range of anomalies without any predefined anomalous patterns or examples to follow. At the same time, this approach is particularly valuable in real-world scenarios, where obtaining labeled data may be costly or impractical. Consequently, employing methods capable of detecting deviations from a purely data-driven inferred notion of normality becomes essential.

To solve the task of unsupervised anomaly detection, various anomaly detection algorithms have been proposed. Methods such as AnoGan [43] aim to model the probability density distribution of normal samples and consider samples in low-density regions as abnormal. Other algorithms, like KNN [14], consider samples that are further away in the variable space from other samples as more anomalous. Approaches like Isolation Forests [39] measure how easily

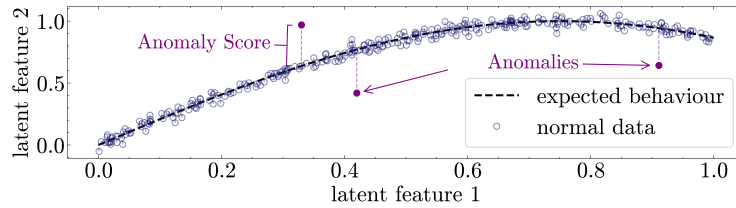


Fig. 1: Example of surrogate anomaly detection: Anomalies are detected by learning a representation that encodes the regular patterns of normal data and measuring deviations from the expected behavior.

samples can be separated. Still, it can be shown that both of these alternative approaches effectively model densities too [7,20].

While density estimation is an intuitive solution to anomaly detection, any method based on this approach has two fundamental drawbacks. First, they do not scale well to high-dimensional data, which is known as the curse of dimensionality [3]. Second, there is usually no explicit modeling of the data-generating process involved, which limits how well the method generalizes to new data.

In contrast to modeling complex density distributions for high dimensional input data, we consider in this paper algorithms that learn low dimensional representations as approximations of the underlying regular patterns in the data. More specifically, we are inspired by surrogate models in engineering applications [33], where simple surrogate models are used to approximate more complex or expensive processes. Similarly, we search for models that capture a pattern of the underlying data-generating process instead of modeling the whole distribution, which we will subsequently refer to as surrogate anomaly detection models. An illustrative toy example of such a model is shown in Figure 1.

Some existing algorithms can be considered instantiations of such surrogate anomaly detection models. Autoencoder [53] and PCA-based anomaly detection [9] learns an identity function to compress regular data and measure the deviation of anomalies as the reconstruction error. DeepSVDD [51] learns a representation in which regular data can be modeled by mapping it to a lower dimensional constant, where anomalies deviate highly from this constant value. Because these algorithms don’t need to model the entire density distribution in its original input space, they scale more effectively to high dimensional data [13].

However, these methods are not without limitations either. Unlike density estimation methods, the objective of training a surrogate is much less well-defined, leading to an unlimited amount of options on how to create such a surrogate, each with its own drawbacks. We exploit this variability to formalize the idea of surrogate models into a blueprint for creating arbitrary surrogates. Within this framework, we propose five axioms that an ideal surrogate model should satisfy. Based on these, we suggest a new surrogate AD algorithm called DEAN, which, to the best of our knowledge, is the first algorithm adhering to all of them.

We evaluate our algorithm by following the procedure outlined in a recent benchmark survey paper [21], comparing it against 19 competitors across 121 datasets. Our algorithm performs highly competitively, showing only minor performance differences with the best non-surrogate competitors and outperforming all other surrogate-based methods.

Our main contributions are: (1) the formulation of a general framework for surrogate anomaly detection; (2) the establishment of guiding axioms for designing optimal surrogate algorithms; and (3) the development and comprehensive evaluation of a novel algorithm based on these principles.

To ensure the reproducibility of our results, our implementation, as well as our appendices containing further details about our experiments, are publicly available at github.com/KDD-OpenSource/DEAN.

2 Related Work

This section reviews related work with a focus on three key aspects: unsupervised anomaly detection, emphasizing approaches that extract meaningful patterns, ensemble methods, which, as discussed later, may enable the extraction of diverse patterns, and surrogate models in a more general context.

2.1 Unsupervised Anomaly Detection

Anomaly detection in an unsupervised setting inherently faces the challenge of defining a suitable objective without ground truth labels. A common suggestion is to model the densities of normal, expected samples [46], under the assumption that samples in low-density regions are less likely to be generated by the same process as normal data, and are therefore more likely to be anomalies. However, density estimation fundamentally suffers from the so called curse of dimensionality [3], which limits how well these algorithms work on high-dimensional data.

Instead of modeling the densities of normal samples, certain anomaly detection methods extract a characteristic pattern that normal samples typically satisfy and test, whether new samples conform to this pattern. Examples of this are DeepSVDD [51], which tries to learn a representation in which every sample is mapped close to a certain point, or reconstruction-based methods like Autoencoder [53] and PCA [9], which try to learn a lower dimensional latent representation, that captures all necessary information to reconstruct (only) normal samples.

These anomaly detection algorithms are less affected by the curse of dimensionality, because latent patterns typically do not increase significantly in complexity with additional features [13]. Still, these algorithms also have flaws. DeepSVDD requires careful training or will simply not perform well [21], and using reconstruction-based algorithms requires choosing a suitable size of the latent space, which is difficult without feedback through labeled anomalies [42]. Thus, in this paper, we generalize such models and suggest an optimal approach based on axioms that outline essential properties.

2.2 Ensemble Methods

Ensembles are a powerful method in machine learning [2]; techniques such as bagging, boosting, and stacking are well-established for combining multiple submodels into a superior model. In supervised tasks, the availability of labels facilitates the coordination of submodels [11]. While there exist approaches to mitigate this, for example with synthetic ground truth anomalies [18], in unsupervised settings, the absence of labels complicates this process. Thus, many unsupervised anomaly detection approaches simply aggregate the anomaly scores produced by various algorithms [30]. This strategy takes advantage of the fact, that errors made by diverse submodels tend not to be repeated across the entire ensemble [8].

One effective approach is the use of homogeneous ensembles, which merge many similar and simple submodels that, although weak individually, collectively yield robust performance through a combination of specialization and diversity [39,12]. Moreover, model-independent methods such as feature bagging [36] can further enhance diversity by ensuring that submodels specialize in different subsets of data dimensions, which can also improve explainability [29].

2.3 Surrogate Models

Surrogate models are simplified abstractions of more complex or computationally expensive models [33]. In engineering, they are commonly employed when, for example, physical simulations are too costly [45]. In machine learning, surrogates have been used to approximate, accelerate, or explain other machine learning models. This has been applied to uncertainty estimation [54], explainability (both globally [44], and locally [50]), surrogate task-based models [57], and to accelerate anomaly detection [16].

In contrast, we propose surrogate anomaly detection to directly learn an approximation of the regular patterns in the input data. This approach enables a more reliable measurement of deviations compared to traditional density estimation methods, particularly for complex, high-dimensional data.

3 Theory of Surrogate Anomaly Detection

In engineering applications, surrogate models are frequently employed when the underlying processes are too complex to model directly. Similarly, when it comes to anomaly detection, it may be impractical to model a complex distribution directly. Here, we define a *surrogate* as a model that approximately learns characteristic patterns of normal samples and identifies anomalies through deviations from these patterns.

This idea can be formalized by requiring that a learnable function $f : \mathbb{R}^d \rightarrow \mathbb{R}^k$ approximates a target pattern $g : \mathbb{R}^d \rightarrow \mathbb{R}^k$ over the set $X \subset \mathbb{R}^d$ of normal data samples:

$$f(x) \approx g(x), \quad \forall x \in X \tag{1}$$

For example, consider a dataset where each normal sample satisfies $x_0 = x_1$. In this case, if we set $f(x) = x_0$ and $g(x) = x_1$, any significant discrepancy between x_0 and x_1 would indicate an anomaly. In practice, f may be realized by training a neural network to map high-dimensional input data to a lower-dimensional latent space (with $k \ll d$), where the underlying structure of normality is more apparent.

The target pattern g may be chosen in various ways. For instance, in an autoencoder g is the identity function, i.e., $g_{AE}(x) = x$. However, since sufficiently expressive neural networks are universal function approximators [41], there are no inherent restrictions on the choice of g ; the only requirement is that it represents a pattern that is largely invariant across normal samples.

To quantify the extent to which a sample x deviates from the learned pattern, we can measure the difference between $f(x)$ and $g(x)$:

$$score(x) = \|f(x) - g(x)\| \quad (2)$$

Since the goal is to ensure that normal samples conform to the learned pattern, we can minimize the aggregate deviation over the training data X_{train} by employing the loss function

$$\mathcal{L} = \sum_{x \in X_{train}} score(x) \quad (3)$$

A critical observation is that while minimizing this loss drives $f(x)$ closer to $g(x)$ for normal samples, it does not directly enforce a high anomaly score for abnormal ones. Consequently, surrogate models may require additional mechanisms to avoid trivial solutions, as discussed in [25].

In summary, Equations 2 and 3 provide a general framework for developing surrogate anomaly detection algorithms. Although any function g consistent with the definition may be used, its effectiveness in yielding a well-performing anomaly detector may vary considerably.

3.1 Surrogate Axioms

To guide the selection of the pattern function g in our surrogate model, we propose five axioms that an optimal surrogate algorithm should satisfy. We assume that a performance measure $m(f)$ (e.g., AUC-ROC) exists, which evaluates how well a model separates anomalies from normal samples.

First, note that the comparison in Equation 2 depends not only on the relative deviation of $f(x)$ from $g(x)$, but also on the magnitude of $\|g(x)\|$. If $\|g(x)\|$ varies significantly across samples, this may unfairly bias their anomaly score assignments.

Axiom 1 (Scale Consistency) *The pattern function g should produce outputs of similar scale for all inputs: $\forall x_1, x_2 \in \mathbb{R}^d$ it holds that $\|g(x_1)\| \approx \|g(x_2)\|$.*

An optimal surrogate must also yield similar results under identical training conditions, ensuring that any observed performance is not a mere artifact of random initialization.

Axiom 2 (Reliable Training Procedure) *When learning to approximate g multiple times under identical training conditions, the variance in performance should be small. For learned instances $f_1, \dots, f_n$ it holds that $\text{Var}(m(f_i)) \leq \delta^2$, where the constant $\delta > 0$ is as small as possible.*

It is also crucial to be robust against trivial solutions – functions that (locally) minimize the loss $\mathcal{L}$, yet have no ability to discern between normal and anomalous samples – since such solutions render the model useless for anomaly detection.

Axiom 3 (Robustness to Trivial Solutions) *There should be no trivial solution f_{trivial} such that $\nabla \mathcal{L}(f_{\text{trivial}}) = 0$ and $f_{\text{trivial}}(x) \approx c$ for all $x \in \mathbb{R}^d$ and some constant $c \in \mathbb{R}^k$.*

Hyperparameter selection poses a significant challenge in anomaly detection [15,59]. Thus, an optimal surrogate should exhibit stability under reasonable variations in hyperparameters, that do not fundamentally alter the model’s methodological design or learning dynamics.

Axiom 4 (Hyperparameter Invariance) *For any two reasonable hyperparameter sets H_A and H_B , let f_{H_A} and f_{H_B} be the corresponding learned models. Then the performance difference should be bounded as $|m(f_{H_A}) - m(f_{H_B})| \leq \eta$, where $\eta > 0$ is chosen to be as small as possible.*

Finally, because anomalies can be both complex and subtle, it is imperative that the surrogate model possesses sufficient expressive power. The model must be able to capture intricate patterns in the data, allowing it to accurately distinguish between normal and anomalous behavior.

Axiom 5 (Complex Pattern Learning) *The learnable function f needs to be represented by a universal function approximator, capable of approximating any continuous function $g : \mathbb{R}^d \rightarrow \mathbb{R}^k$ to arbitrary precision on subsets of $\mathbb{R}^d$.*

3.2 Axiom Compliance

Both of the most established deep anomaly detection paradigms that conform to our surrogate definition exhibit significant deviations from the proposed axioms, highlighting inherent limitations in their design.

Autoencoder: An Autoencoder [53] defines its surrogate usually via the identity function $g_{AE}(x) = x$, training a neural network to reconstruct its input while enforcing a compression step to prevent the trivial layer-by-layer identity mapping. However, this approach violates Axiom 4 because the latent dimensionality must be carefully chosen, which critically affects performance. Moreover, autoencoders can converge to local minima – such as outputting the mean of

the training samples (violating Axiom 3) – and the lack of consistent scaling in g results in biased anomaly scores (violating Axiom 1).

DeepSVDD: DeepSVDD [51] constructs its surrogate model using a constant pattern function $g_{SVDD}(x) = c$, where c is a predetermined constant, usually chosen as the mean output of the initialized network. To mitigate the risk of learning a trivial constant prediction, the method suggests avoiding bounded activation functions and removing the learnable shifts⁴ from each network layer. Unfortunately, the latter restriction limits the network’s capacity to learn complex patterns, thereby breaching either Axiom 3 or Axiom 5.

4 A Minimal Surrogate: The DEAN Model

Given that no surrogate known to us adheres to all aforementioned axioms, we propose a novel deep learning-based approach. We observe that increased complexity in the pattern function g often leads to arbitrary weighting of different samples (Axiom 1) and intensifies challenges during training for the function f that needs to be learned (Axioms 2 and 3). Thus, we advocate for selecting the simplest possible function g that adequately identifies the essential data patterns.

Depending on the measure of complexity, one might consider $g_0(x) = 0$ as the simplest option. However, this surrogate violates Axiom 3 by introducing a local minimum where every weight in the last layer becomes zero [51]. Although such a minimum might not always be reached [27] or could be avoided using regularization, these strategies would, in turn, compromise our remaining axioms. Instead, we propose $g_{DEAN}(x) = 1$ and the surrogate generated by it. While a local minimum may still occur via the learnable shifts in the final layer, it is more manageable, as we demonstrate later (Section 4.2).

Thus, we train a neural network to output a constant value of 1. In accordance with our framework, the loss and score functions are defined as

$$\mathcal{L} = \sum_{x \in X_{train}} \|f(x) - 1\|, \quad score(x) = \|f(x) - q\| \quad (4)$$

where

$$q = \frac{1}{\|X_{train}\|} \sum_{x_T \in X_{train}} f(x_T) \approx 1$$

This choice of q ensures that the distribution of normal samples is centered, thereby improving the robustness of the anomaly score.

Since we only learn a one-dimensional pattern with this approach, the model may, in the worst case, fix only one feature as a function of the remaining ones, which can lead to an increased number of false negatives. While one might extend g to a higher-dimensional constant vector $g(x) = (1, 1, \dots, 1)^T$, this typically results in significantly correlated outputs across the network’s features and may violate Axiom 4. To address these challenges, we propose using an ensemble of surrogates. Specifically, we combine many independently trained submodels with

⁴ We refer to the bias term of a neural network as "learnable shift" to reduce confusion.

g_{DEAN} into a more effective model F . An integer constant, denoted by $power$, guides the aggregation of these models:

$$score_F(x) = \frac{1}{\|F\|} \sum_{f_i \in F} \|f_i(x) - q_i\|^{power} \quad (5)$$

where each q_i is computed analogously to q . Owing to the simplicity of our surrogate, training each neural network takes only seconds, thus allowing us to combine a large number of submodels. The use of fully independent networks facilitates parallelization, reduces the correlation among learned patterns, and permits the application of ensemble methods such as feature bagging [36] to further improve diversity and runtime consistency in high-dimensional settings. Feature bagging additionally can ensure a constant number of features for each submodel, resulting in a close-to-constant runtime for higher dimensional data.

We refer to this overall setup as *DEAN* (Deep Ensemble ANomaly detection).

4.1 DEAN Parameterization

In our instantiation of DEAN, we advocate for a diverse ensemble of simple and efficient feedforward networks.

Network Architecture: We use a basic Multi-Layer Perceptron (MLP) with only a few hidden layers. Hidden layers are constructed with bias terms and use ReLU activations, while the output layer excludes a bias term and employs a SELU activation to mitigate the risk of dead neurons.

Ensemble Structure: A large ensemble size allows specialization in a variety of different patterns. For high-dimensional datasets, feature bagging is used to promote the diversity of submodels by training each on a random subset of the available features. For datasets with few features, all of them may be used to ensure critical correlations are captured.

Power Parameter: The *power* parameter allows controlling the sensitivity of the aggregated anomaly score. A higher value accentuates significant deviations in one model over multiple smaller deviations across models.

Training Configuration: A lower-than-standard learning rate is paired with a relatively high batch size. This configuration not only stabilizes training but also encourages the network to converge towards local minima, which is beneficial for ensemble diversity.

4.2 Axiom Compliance of DEAN

DEAN is designed to fulfill all of our surrogate axioms. Since $g(x) = 1$ for all x , Axiom 1 (Scale Consistency) is satisfied. The compliance with Axioms 2 (Reliable Training Procedure) and 4 (Hyperparameter Invariance) is best evaluated via experimental comparisons (see Section 5.4). We now discuss the more challenging Axioms 3 (Robustness to Trivial Solutions) and 5 (Complex Pattern Learning). These are non-trivial because the meaningless function $f_{trivial}(x) = 0 \cdot x + 1$ perfectly minimizes the DEAN loss and is independent of its input.

Axiom 3: Given its ensemble structure, DEAN inherently mitigates trivial solutions. In the event that they occur, their contribution to the ensemble anomaly score is zero since $f_{trivial}(x) = 1$ for all x implies that $\|f_{trivial}(x) - q\| = 0$. Thus, with a sufficient number of submodels, the overall performance of DEAN remains unaffected by any individual trivial solution.

Axiom 5: DeepSVDD addresses a similar issue of trivial solutions by removing learnable shifts entirely [51]; however, this approach limits the network’s capacity and violates Axiom 5. To still mitigate this risk, while preserving expressiveness, we remove learnable shifts only from the final layer. This adjustment increases the complexity required to achieve a trivial solution, making it less likely to be reached during training, while ensuring that the network retains the ability to approximate any function (see Appendix B).

5 Experimental Evaluation

To experimentally evaluate our method, we refer to the protocol outlined in the survey paper ADBench [21]. ADBench recommends 121 datasets (57 of which are entirely uncorrelated) for benchmarking unsupervised anomaly detection algorithms, as well as a set of baseline algorithms to compare against.

5.1 Experimental Setup

Following the approach of ADBench, we compare DEAN against a total of 19 state-of-the-art algorithms. This includes KNN [14], LOF [6], CBLOF [22], Isolation Forest [39], PCA [9], DeepSVDD [51], OCSVM [5], LODA [47], HBOS [19], COPOD [37], ECOD [38], SOD [34] and DAGMM [60]. In addition, we consider a regular Autoencoder [53], as it is also a surrogate algorithm, as well as a variational autoencoder [28] and a normalizing flow [49] as deep learning density-based competitors. To further capture recent advances not originally considered in ADBench, we also compare against NeuTral-AD [48] (which leverages contrastive learning), DTE [40] (based on diffusion models), and GOAD [4] (which employs geometric transformations). In contrast to ADBench, all models are trained on uncontaminated data (one-class setting).

For the parameterization of DEAN, we adhere to the guidelines proposed in Section 4.1 in order to train an ensemble of 100 submodels for 50 epochs each, using early stopping with a patience of 10 epochs. We adopt the following specific hyperparameter choices: A feedforward neural network with three hidden layers of 255 neurons each. A lower-than-standard learning rate of 0.0001 and a rather high batch size of 512. Feature bagging with 200 random features per model for datasets containing at least 200 features. A power parameter set to 9, emphasizing pronounced deviations in anomaly detection. We consider variations in ensemble size and other hyperparameters in Sections 5.3 and 5.4. Detailed information regarding the implementation and parameterization of the compared methods can be found in our code repository.

5.2 Anomaly Detection Performance

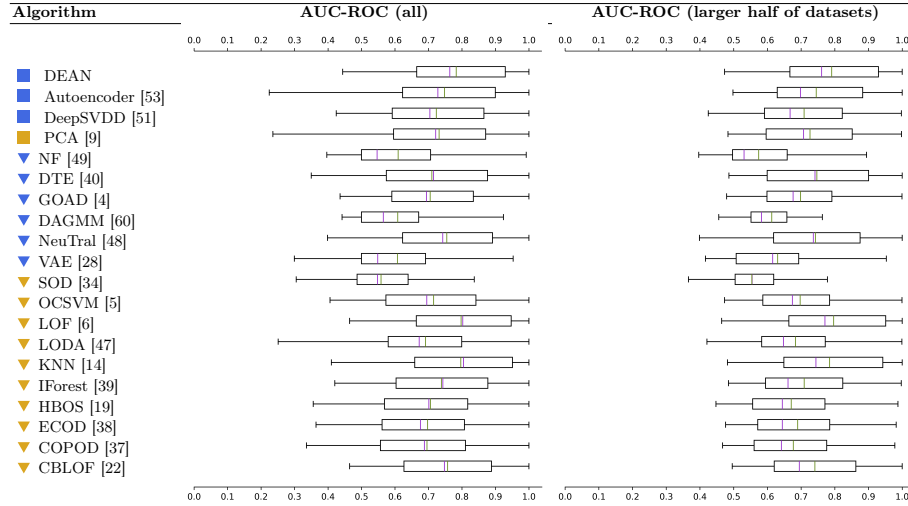


Table 1: Distribution of AUC-ROC performance for all evaluated algorithms. Deep learning models (blue) and shallow models (yellow) are differentiated by surrogate status (squares for surrogates, triangles for non-surrogates). Mean and median values are shown in green and purple, respectively.

Our primary performance evaluation metric is the Area Under the Receiver Operating Characteristic (AUC-ROC). For additional insight, we provide a complementary evaluation based on the Area Under the Precision-Recall Curve (AUC-PR), alongside individual results for each dataset, in Appendix E and F.

A summarization of the observed AUC-ROC performance is given in Table 1. Consistent with the results from ADBench, our analysis shows that no method outperforms all others in a statistically significant manner across all datasets. Our algorithm performs highly competitively when averaged across all datasets and is only slightly outperformed by KNN and LOF. These competitors do not scale well to the more challenging datasets. Thus, when only considering the larger half of the datasets studied here, the median performance of DEAN is higher than all competitors considered.

To further illustrate our findings and provide a concise ranking of performance, we provide the critical difference diagrams in Figure 2. In these diagrams, a Friedman test [17] is used to determine if significant differences exist between algorithm performances (measured in AUC-ROC), and algorithms with no significant differences are connected using a Wilcoxon test [56]. We consider p -values below $p \leq 5\%$ after Bonferroni-Holm correction [26] to be significant.

Notably, DEAN performs significantly stronger than every other surrogate or deep learning algorithm, with the exception of NeuTral (see Figure 2a). Similar to

ADBench, widely recognized shallow algorithms such as KNN, LOF, and CBLOF remain strong competitors. Our algorithm outperforms CBLOF, but does not quite achieve the same average rank as KNN and LOF. We attribute this outcome to the fact that benchmark datasets are often low-dimensional, contain a large number of samples, and exhibit anomalies that are relatively simple in nature. This combination favors distance-based methods, as differences in local densities are more pronounced; however, they may not capture the complexity encountered in real-world applications. Moreover, the lazy learning paradigm intrinsic to KNN and LOF, which necessitates the retention of training instances during inference, is well known to scale badly to large, high-dimensional datasets [1].

For instance, when considering only the larger half of the datasets, DEAN emerges as the best fit, outperforming every competitor (see Figure 2b). This performance advantage, coupled with its scalability and robust learning framework, makes DEAN particularly well suited for advanced tasks where the expressive power of neural networks is needed without introducing unnecessary complexity.

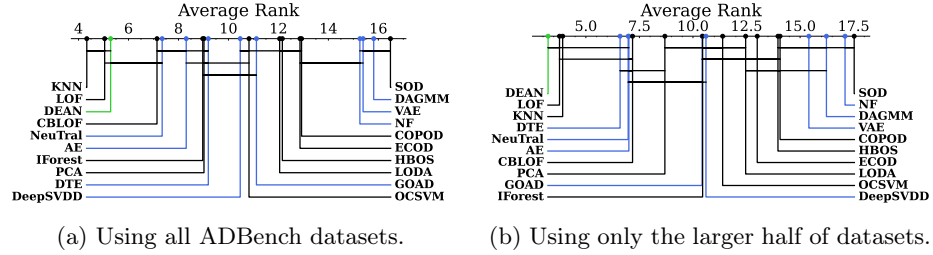


Fig. 2: Critical difference diagrams comparing the AUC-ROC performance. A lower rank indicates better performance, while algorithms with no statistically significant differences are connected by a horizontal line. DEAN is depicted in green, other deep learning algorithms in blue.

5.3 Runtime and Ensemble Analysis

Figure 3 provides an overview of our runtime measurements and the impact of using (larger) ensembles on the performance of deep learning-based surrogate models. To this end, Figure 3a reports both the median and maximum runtimes across datasets to account for the substantial variability in dataset sizes. All experiments were conducted on a system running Ubuntu 22.04.3 LTS, powered by an Intel® Xeon® w9-3495X processor with a base clock of approximately 3400 MHz and turbo boost frequencies up to around 4500 MHz and with 495 GB of RAM available. The runtime measurements were obtained using single-core execution for each algorithm to ensure a fair comparison.

As expected, deep learning methods generally exhibit longer runtimes, with the DEAN ensemble showing a median runtime of approximately 15 minutes

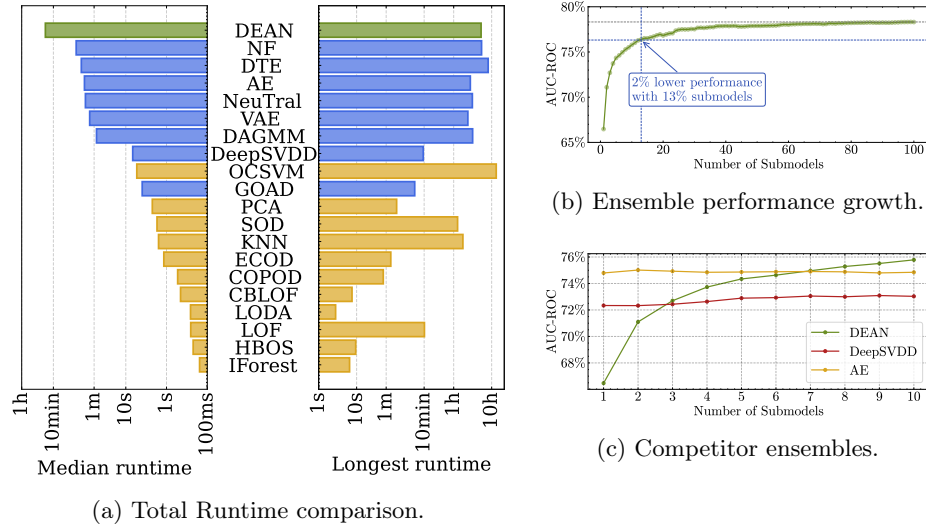


Fig. 3: (a) Overall runtime overview across all datasets; DEAN is depicted in green, other deep learning algorithms in blue, and shallow algorithms in yellow. (b) and (c) average AUC-ROC performance changes with varying ensemble size.

per dataset. Notably, for a DEAN ensemble comprising 100 submodels, this corresponds to an average training time of less than 9 seconds per submodel. However, it is important to emphasize that deep learning methods are particularly well-suited for GPU acceleration, and DEAN, as an ensemble method of independently optimized submodels, can be almost perfectly parallelized. Furthermore, due to the use of feature bagging, the worst-case runtime scenario is significantly mitigated, with most deep learning approaches requiring comparable or even longer runtimes.

Naturally, the runtime is also rather sensitive to the number of submodels. As illustrated in Figure 3b, while increasing the number of submodels improves performance, the relationship is non-linear. Using only 13 submodels results in an average performance that is merely 2% lower than that achieved with 100 submodels, yet it requires approximately 87% less training time. At the same time, the continued performance improvement with additional submodels reflects the high variance of the individual models used in DEAN, incentivized by the simplicity of the submodels. In contrast, Figure 3c shows that ensembles based on Autoencoder or DeepSVDD methods exhibit nearly constant performance, likely due to their complex, less diverse submodel characteristics.

5.4 Evaluation of Axiom Compliance

Compliance with Axioms 2 and 4 is difficult to assess theoretically, therefore we evaluate these properties experimentally on the same datasets. For Axiom 2,

Figure 4 demonstrates that the repetition uncertainty of DEAN – calculated as the standard deviation across 10 runs per datasets and then averaged – is lower than that observed for other surrogate deep learning algorithms under identical training conditions. For Axiom 4, we present DEAN’s performance when evaluated with modified hyperparameter sets. Since the average performances are nearly equal, one may argue that the influence of such modifications is negligible. This is also in stark contrast to DeepSVDD [21] and Autoencoder [35] behavior.

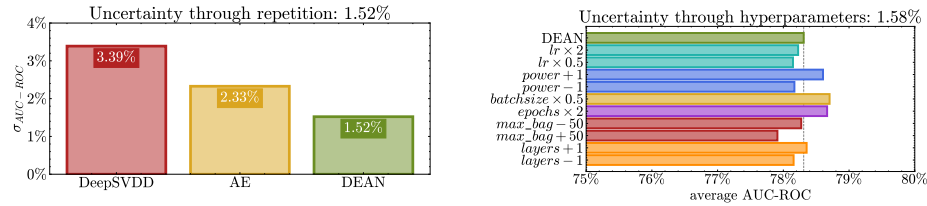


Fig. 4: Left: Repetition uncertainty for various surrogate algorithms, Right: DEAN performance with varied hyperparameters and the resulting uncertainty.

6 Anomaly Detection Beyond Benchmarks

While DEAN reliably achieves competitive performance with an easy-to-configure parameterization, real-world applications often demand flexibility beyond standard benchmarks. The simplicity of our submodels and the inherent ensemble structure render DEAN highly adaptable.

For instance, the ensemble structure facilitates explainability via Shapley values [29]; feature bagging helps mitigate the high computational cost of such methods. In addition, the ensemble character natively supports a distributed implementation through federated learning [55]. The ensemble also enables the incorporation of secondary requirements, such as robustness against adversarial attacks by pruning or reweighting less robust submodels [8]. The simplicity of each submodel also permits modifications in the training procedure to incorporate additional information [31], such as in semi-supervised anomaly detection [52] or outlier exposure [23]. Moreover, employing different machine learning models within the DEAN framework can yield lightweight variants suitable for resource-constrained devices such as IoT systems [32].

To summarize, we see three major ways DEAN can be adapted: (1) altering the selection of submodels, (2) adjusting the ensemble weighting, and (3) modifying the submodel training procedure. As a proof-of-concept, we apply all three adaptations for the task of fair anomaly detection [58] (see Appendix B).

7 Conclusion

In this paper, we present the first systematic study of surrogate models for anomaly detection and establish a comprehensive framework for constructing such models from mathematical functions. We propose five axioms that any optimal surrogate anomaly detection algorithm should satisfy and employ these axioms to develop DEAN, a novel algorithm that meets all of them.

An extensive evaluation demonstrates that it not only performs competitively – particularly excelling over other deep learning-based methods and alternative surrogates – but also offers exceptional adaptability. In the future, we believe that the axiomatic design of DEAN, based on an ensemble of simple submodels, can furthermore facilitate straightforward modifications to enhance secondary anomaly detection goals, like explainability, adversarial robustness, or fairness.

Acknowledgements

This work was supported by the Lamarr-Institute for ML and AI, the Research Center Trustworthy Data Science and Security, the Federal Ministry of Education and Research of Germany and the German federal state of NRW. The Linux HPC cluster at TU Dortmund University, a project of the German Research Foundation, provided the computing power.

References

1. Aggarwal, C.C.: Outlier Analysis. Springer (2013), <http://dx.doi.org/10.1007/978-1-4614-6396-2>
2. Aggarwal, C.C., Sathe, S.: Theoretical foundations and algorithms for outlier ensembles. SIGKDD Explor. Newsl. **17**(1), 24–47 (sep 2015). <https://doi.org/10.1145/2830544.2830549>
3. Bellman, R.: Dynamic programming. Science **153**(3731), 34–37 (1966)
4. Bergman, L., Hoshen, Y.: Classification-based anomaly detection for general data. In: International Conference on Learning Representations (ICLR) 2020 (2020), https://openreview.net/forum?id=H1lK_lBtvS
5. Bounsiar, A., Madden, M.G.: One-class support vector machines revisited. In: 2014 International Conference on Information Science Applications (ICISA). pp. 1–4 (2014). <https://doi.org/10.1109/ICISA.2014.6847442>
6. Breunig, M., Kröger, P., Ng, R., Sander, J.: Lof: Identifying density-based local outliers. vol. 29, pp. 93–104 (06 2000). <https://doi.org/10.1145/342009.335388>
7. Buschjäger, S., Honysz, P.J., Morik, K.: Randomized outlier detection with trees. International Journal of Data Science and Analytics **13**(2), 91–104 (Mar 2022). <https://doi.org/10.1007/s41060-020-00238-w>
8. Böing, B., Klüttermann, S., Müller, E.: Post-robustifying deep anomaly detection ensembles by model selection. In: 2022 IEEE International Conference on Data Mining (ICDM). pp. 861–866 (2022). <https://doi.org/10.1109/ICDM54844.2022.00098>

9. Callegari, C., Gazzarrini, L., Giordano, S., Pagano, M., Pepe, T.: A novel pca-based network anomaly detection. pp. 1 – 5 (07 2011). <https://doi.org/10.1109/icc.2011.5962595>
10. Castro, C.M.R.I.I.P.M.: Detecting falls as novelties in acceleration patterns acquired with smartphones. PLoS ONE **9**, None (2014). <https://doi.org/10.1371/journal.pone.0094811>
11. Chakraborty, D., Narayanan, V., Ghosh, A.: Integration of deep feature extraction and ensemble learning for outlier detection. Pattern Recognition **89**, 161–171 (2019). <https://doi.org/https://doi.org/10.1016/j.patcog.2019.01.002>
12. Chen, J., Sathe, S., Aggarwal, C., Turaga, D.: Outlier Detection with Autoencoder Ensembles, pp. 90–98 (06 2017). <https://doi.org/10.1137/1.9781611974973.11>
13. Chen, W., Li, H., Li, J., Arshad, A.: Autoencoder-based outlier detection for sparse, high dimensional data. pp. 2735–2742 (12 2020). <https://doi.org/10.1109/BigData50022.2020.9378325>
14. Cunningham, P., Delany, S.: k-nearest neighbour classifiers. Mult Classif Syst **54** (04 2007). <https://doi.org/10.1145/3459665>
15. Ding, X., Zhao, L., Akoglu, L.: Hyperparameter sensitivity in deep outlier detection: Analysis and a scalable hyper-ensemble solution. In: Advances in Neural Information Processing Systems. vol. 35, pp. 9603–9616. Curran Associates, Inc. (2022), https://proceedings.neurips.cc/paper_files/paper/2022/file/3e9113e2bc2e700baa7d765470f140e1-Paper-Conference.pdf
16. Flusser, M., Somol, P.: Efficient anomaly detection through surrogate neural networks. Neural Computing and Applications **34**, 1–15 (06 2022). <https://doi.org/10.1007/s00521-022-07506-9>
17. Friedman, M.: A comparison of alternative tests of significance for the problem of \$m\$ rankings. Annals of Mathematical Statistics **11**, 86–92 (1940)
18. Fung, C., Qiu, C., Li, A., Rudolph, M.: Model selection of anomaly detectors in the absence of labeled validation data. IEEE Transactions on Artificial Intelligence (2025). <https://doi.org/10.1109/TAI.2025.3562505>
19. Goldstein, M., Dengel, A.R.: Histogram-based outlier score (hbos): A fast unsupervised anomaly detection algorithm (2012), <https://api.semanticscholar.org/CorpusID:3590788>
20. Gu, X., Akoglu, L., Rinaldo, A.: Statistical analysis of nearest neighbor methods for anomaly detection. In: Neural Information Processing Systems (NeurIPS). pp. 10921–10931 (2019), <http://dblp.uni-trier.de/db/conf/nips/nips2019.html#GuAR19>
21. Han, S., Hu, X., Huang, H., Jiang, M., Zhao, Y.: Adbench: Anomaly detection benchmark. In: Neural Information Processing Systems (NeurIPS) (2022), https://proceedings.neurips.cc/paper_files/paper/2022/file/cf93972b116ca5268827d575f2cc226b-Paper-Datasets_and_Benchmarks.pdf
22. He, Z., Xu, X., Deng, S.: Discovering cluster-based local outliers. Pattern Recognition Letters **24**(9), 1641–1650 (2003). [https://doi.org/https://doi.org/10.1016/S0167-8655\(03\)00003-5](https://doi.org/https://doi.org/10.1016/S0167-8655(03)00003-5)
23. Hendrycks, D., Mazeika, M., Dietterich, T.: Deep anomaly detection with outlier exposure (2019), <https://arxiv.org/abs/1812.04606>
24. Hilal, W., Gadsden, S.A., Yawney, J.: Financial fraud: a review of anomaly detection techniques and recent advances. Expert systems With applications **193** (2022). <https://doi.org/https://doi.org/10.1016/j.eswa.2021.116429>
25. Hoffer, E., Ailon, N.: Deep metric learning using triplet network. In: Similarity-Based Pattern Recognition (2015), <https://api.semanticscholar.org/CorpusID:2784676>

26. Holm, S.: A simple sequentially rejective multiple test procedure. *Scandinavian Journal of Statistics* **6**(2), 65–70 (1979), <http://www.jstor.org/stable/4615733>
27. Karr, N., Nachman, B., Shih, D.: One-class dense networks for anomaly detection. In: *Proceedings of the Machine Learning and the Physical Sciences Workshop at NeurIPS 2022* (December 2022), https://ml4physicalsciences.github.io/2022/files/NeurIPS_ML4PS_2022_130.pdf
28. Kingma, D., Welling, M.: Auto-encoding variational bayes (2014). <https://doi.org/10.61603/ceas.v2i1.33>
29. Klüttermann, S., Balestra, C., Müller, E.: On the efficient explanation of outlier detection ensembles through shapley values. In: *Advances in Knowledge Discovery and Data Mining*. pp. 43–55 (2024). https://doi.org/10.1007/978-981-97-2259-4_4
30. Klüttermann, S., Müller, E.: Evaluating and comparing heterogeneous ensemble methods for unsupervised anomaly detection (2023). <https://doi.org/10.1109/IJCNN54540.2023.10191405>
31. Klüttermann, S., Müller, E.: About test-time training for outlier detection (2024), <https://arxiv.org/abs/2404.03495>
32. Klüttermann, S., Peka, V., Doeblér, P., Müller, E.: Towards highly efficient anomaly detection for predictive maintenance. In: *2024 International Conference on Machine Learning and Applications (ICMLA)*. pp. 1691–1696 (2024). <https://doi.org/10.1109/ICMLA61862.2024.00261>
33. Koziel, S., Ciaurri, D.E., Leifsson, L.: *Surrogate-Based Methods*, pp. 33–59. Springer Berlin Heidelberg, Berlin, Heidelberg (2011). https://doi.org/10.1007/978-3-642-20859-1_3
34. Kriegel, H.P., Kröger, P., Schubert, E., Zimek, A.: Outlier detection in axis-parallel subspaces of high dimensional data. In: *Advances in Knowledge Discovery and Data Mining*. pp. 831–838 (2009). https://doi.org/10.1007/978-3-642-01307-2_86
35. Kumar, V., Srivastava, V., Mahjabin, S., Pal, A., Klüttermann, S., Müller, E.: Autoencoder optimization for anomaly detection: A comparative study with shallow algorithms. In: *International Joint Conference on Neural Networks (IJCNN)*. <https://psorus.github.io/papers/vikas.pdf>
36. Lazarevic, A., Kumar, V.: Feature bagging for outlier detection. In: *Proceedings of the Eleventh ACM SIGKDD International Conference on Knowledge Discovery in Data Mining*. p. 157–166. KDD '05 (2005). <https://doi.org/10.1145/1081870.1081891>
37. Li, Z., Zhao, Y., Botta, N., Ionescu, C., Hu, X.: Copod: Copula-based outlier detection. In: *2020 IEEE International Conference on Data Mining (ICDM)* (2020). <https://doi.org/10.1109/ICDM50108.2020.00135>
38. Li, Z., Zhao, Y., Hu, X., Botta, N., Ionescu, C., Chen, G.H.: Ecod: Unsupervised outlier detection using empirical cumulative distribution functions. *IEEE Transactions on Knowledge and Data Engineering* **35**(12), 12181–12193 (2023). <https://doi.org/10.1109/TKDE.2022.3159580>
39. Liu, F.T., Ting, K.M., Zhou, Z.H.: Isolation forest. In: *International Conference on Data Mining (ICDM)* (2008). <https://doi.org/10.1109/ICDM.2008.17>
40. Livernoche, V., Jain, V., Hezaveh, Y., Ravanbakhsh, S.: On diffusion modeling for anomaly detection. In: *International Conference on Learning Representations (ICLR) 2024* (2024), <https://openreview.net/forum?id=IR3rk7ysXz>
41. Lu, Z., Pu, H., Wang, F., Hu, Z., Wang, L.: The expressive power of neural networks: a view from the width. In: *Proceedings of the 31st International Conference on Neural Information Processing Systems*. p. 6232–6240. NIPS'17, Curran Associates Inc., Red Hook, NY, USA (2017). <https://doi.org/10.5555/3295222.3295371>

42. Ma, M.Q., Zhao, Y., Zhang, X., Akoglu, L.: The need for unsupervised outlier model selection: A review and evaluation of internal evaluation strategies. *ACM SIGKDD Explorations Newsletter* **25**(1) (2023). <https://doi.org/10.1145/3606274.3606277>
43. Mattia, F.D., Galeone, P., Simoni, M.D., Ghelfi, E.: A survey on gans for anomaly detection (2021), <http://arxiv.org/abs/1906.11632>
44. Monteiro, W.R., Reynoso-Meza, G.: A multi-objective optimization design to generate surrogate machine learning models in explainable artificial intelligence applications. *EURO Journal on Decision Processes* **11**, 100040 (2023). <https://doi.org/10.1016/j.ejdp.2023.100040>
45. Mora-Mariano, D., Flores-Tlacuahuac, A.: A machine learning approach for the surrogate modeling of uncertain distributed process engineering models. *Chemical Engineering Research and Design* **186**, 433–450 (2022). <https://doi.org/10.1016/j.cherd.2022.07.050>
46. Olteanu, M., Rossi, F., Yger, F.: Meta-survey on outlier and anomaly detection. *Neurocomputing* **555**, 126634 (2023). <https://doi.org/10.1016/j.neucom.2023.126634>
47. Pevný, T.: Loda: Lightweight on-line detector of anomalies. *Mach. Learn.* **102**(2) (feb 2016). <https://doi.org/10.1007/s10994-015-5521-0>
48. Qiu, C., Pfrommer, T., Kloft, M., Mandt, S., Rudolph, M.: Neural transformation learning for deep anomaly detection beyond images. In: *International conference on machine learning*. pp. 8703–8714. PMLR (2021), <https://proceedings.mlr.press/v139/qiu21a.html>
49. Rezende, D.J., Mohamed, S.: Variational inference with normalizing flows. In: *Proceedings of the 32nd International Conference on International Conference on Machine Learning - Volume 37*. p. 1530–1538. ICML'15, JMLR.org (2015), <https://dl.acm.org/doi/10.5555/3045118.3045281>
50. Ribeiro, M.T., Singh, S., Guestrin, C.: "why should i trust you?": Explaining the predictions of any classifier. In: *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. p. 1135–1144 (2016). <https://doi.org/10.1145/2939672.2939778>
51. Ruff, L., Vandermeulen, R., Goernitz, N., Deecke, L., Siddiqui, S.A., Binder, A., Müller, E., Kloft, M.: Deep one-class classification. In: *Proceedings of the 35th International Conference on Machine Learning*. pp. 4393–4402 (2018), <https://proceedings.mlr.press/v80/ruff18a.html>
52. Ruff, L., Vandermeulen, R.A., Görnitz, N., Binder, A., Müller, E., Müller, K., Kloft, M.: Deep semi-supervised anomaly detection. In: *8th International Conference on Learning Representations, ICLR 2020, Addis Abeba, Ethiopia, April 26-30, 2020*. OpenReview.net (2020), <https://openreview.net/forum?id=HkgH0TEYwH>
53. Sakurada, M., Yairi, T.: Anomaly detection using autoencoders with nonlinear dimensionality reduction. In: *Proceedings of the MLSDA 2014 2nd Workshop on Machine Learning for Sensory Data Analysis*. p. 4–11. MLSDA'14 (2014). <https://doi.org/10.1145/2689746.2689747>
54. Sudret, B., Marelli, S., Wiart, J.: Surrogate models for uncertainty quantification: An overview. In: *2017 11th European Conference on Antennas and Propagation (2017)*. <https://doi.org/10.23919/EuCAP.2017.7928679>
55. Wang, X., Wang, Y., Javaheri, Z., Almutairi, L., Moghadamnejad, N., Younes, O.S.: Federated deep learning for anomaly detection in the internet of things. *Computers and Electrical Engineering* **108**, 108651 (2023). <https://doi.org/10.1016/j.compeleceng.2023.108651>

56. Wilcoxon, F.: Individual Comparisons by Ranking Methods, pp. 196–202. Springer New York, New York, NY (1992). https://doi.org/10.1007/978-1-4612-4380-9_16
57. Ye, F., Zheng, H., Huang, C., Zhang, Y.: Deep unsupervised image anomaly detection: An information theoretic framework. In: 2021 IEEE International Conference on Image Processing (ICIP). pp. 1609–1613 (2021). <https://doi.org/10.1109/ICIP42928.2021.9506079>
58. Zhang, H., Davidson, I.: Towards fair deep anomaly detection. In: Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency. p. 138–148. FAccT ’21, Association for Computing Machinery, New York, NY, USA (2021). <https://doi.org/10.1145/3442188.3445878>
59. Zhao, Y., Rossi, R., Akoglu, L.: Automatic unsupervised outlier model selection. In: Advances in Neural Information Processing Systems. vol. 34, pp. 4489–4502 (2021), https://proceedings.neurips.cc/paper_files/paper/2021/file/23c894276a2c5a16470e6a31f4618d73-Paper.pdf
60. Zong, B., Song, Q., Min, M.R., Cheng, W., Lumezanu, C., ki Cho, D., Chen, H.: Deep autoencoding gaussian mixture model for unsupervised anomaly detection. In: International Conference on Learning Representations (2018), <https://api.semanticscholar.org/CorpusID:51805340>