



A disk failure prediction model for multiple issues*

Yunchuan GUAN^{†1}, Yu LIU^{†‡2}, Ke ZHOU¹, Qiang LI³, Tuanjie WANG³, Hui LI³

¹Wuhan National Laboratory for Optoelectronics, Huazhong University of Science and Technology, Wuhan 430074, China

²School of Computer Science and Technology, Huazhong University of Science and Technology, Wuhan 430074, China

³Inspur Electronic Information Industry Co., Ltd., Beijing 250000, China

[†]E-mail: hustgyc@hust.edu.cn; liu_yu@hust.edu.cn

Received Oct. 19, 2022; Revision accepted June 13, 2023; Crosschecked July 1, 2023

Abstract: Disk failure prediction methods have been useful in handling a single issue, e.g., heterogeneous disks, model aging, and minority samples. However, because these issues often exist simultaneously, prediction models that can handle only one will result in prediction bias in reality. Existing disk failure prediction methods simply fuse various models, lacking discussion of training data preparation and learning patterns when facing multiple issues, although the solutions to different issues often conflict with each other. As a result, we first explore the training data preparation for multiple issues via a data partitioning pattern, i.e., our proposed multi-property data partitioning (MDP). Then, we consider learning with the partitioned data for multiple issues as learning multiple tasks, and introduce the model-agnostic meta-learning (MAML) framework to achieve the learning. Based on these improvements, we propose a novel disk failure prediction model named MDP-MAML. MDP addresses the challenges of uneven partitioning and difficulty in partitioning by time, and MAML addresses the challenge of learning with multiple domains and minor samples for multiple issues. In addition, MDP-MAML can assimilate emerging issues for learning and prediction. On the datasets reported by two real-world data centers, compared to state-of-the-art methods, MDP-MAML can improve the area under the curve (AUC) and false detection rate (FDR) from 0.85 to 0.89 and from 0.85 to 0.91, respectively, while reducing the false alarm rate (FAR) from 4.88% to 2.85%.

Key words: Storage system reliability; Disk failure prediction; Self-monitoring analysis and reporting technology (SMART); Machine learning

<https://doi.org/10.1631/FITEE.2200488>

CLC number: TN333

1 Introduction

Hard disk failures are responsible for most storage systems' failures, resulting in the need to improve disk reliability (Pinheiro et al., 2007). Machine learning (ML) based models that learn with log data, e.g., self-monitoring analysis and reporting technology (SMART), are highly accurate in disk failure prediction, and they deal with issues by distinguishing the corresponding properties. This pattern is

good at learning about a specific issue, but tricky at solving multiple ones. As shown in Fig. 1, the performance of the prediction model degrades in cases of heterogeneous disks (Rincón et al., 2017), model aging (Xiao et al., 2018), and minority samples (Zhang J et al., 2020b). In addition, the physical location of the disk and its relationship to the server cluster can be correlated with disk failures (Lu et al., 2020; Han et al., 2021; Luo et al., 2021). Such environment-related factors may further affect the performance of the failure prediction model, resulting in environmental variation issues. Conventional methods consider the disk model to be key in distinguishing sample categories for the heterogeneous disks issue. Then, they use the samples in the same category for

[‡] Corresponding author

* Project supported by the National Natural Science Foundation of China (No. 61902135) and the Shandong Provincial Natural Science Foundation, China (No. ZR2019LZH003)

ORCID: Yunchuan GUAN, <https://orcid.org/0000-0002-2619-4312>; Yu LIU, <https://orcid.org/0000-0002-1964-9278>

© Zhejiang University Press 2023

training. Similarly, the key is the time period for the model aging issue. However, because the above issues often exist simultaneously, selecting the key and preparing the training data become inescapable challenges.

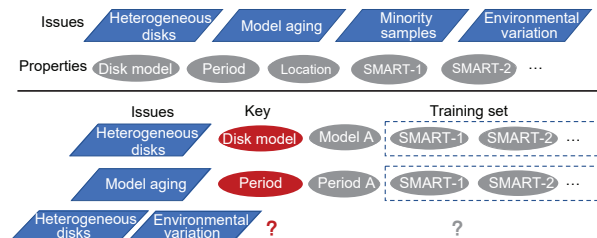


Fig. 1 The problem of training data preparation for multi-issue learning

Review of the existing models that deal with multiple issues shows that they all avoid the above challenges and use the model's ability to handle multiple issues. For example, Xie et al. (2018) built predictors for different issues and tried to use a cross-validation method to vote for the most precise predictor. Zhang J et al. (2020a) solved the minority samples and heterogeneous disks issues by learning samples in the Siamese network framework, which is sensitive to differences in disk models. Although these methods provide proven solutions over multiple issues, they ignore the conflicts between predictions for different issues (see Sections 3 and 6.4), which creates sub-optimal results.

As a result, we believe that learning a disk failure prediction model for multiple issues requires to address the following two challenges:

Generalizability: It is essential to design a framework for multiple issues. However, differences in dataset specifications can challenge the generalization ability of the failure prediction method. For example, the disks in the Backblaze data center (Backblaze, 2018) have the issues of model aging, heterogeneous disks, and minority samples, while the disks in the dataset open-sourced by Lu et al. (2020) have the issues of heterogeneous disks and environmental variation.

Uniformity: Addressing various issues in a unified framework is challenging because the solutions to different issues often conflict with one another, e.g., the issues of model aging and minority samples. The former needs samples from the latest period to perform online updating (Xiao et al., 2018), whereas the latter requires a large number of samples from

other disk models or from a long period to achieve transfer learning (Zhang J et al., 2020b). To reconcile similar conflicts, it is necessary to find the commonality among these issues.

To address these challenges, we analyze the relationships between the issues of heterogeneous disks, model aging, and environmental variation, separately. We find that they can be interpreted as the same data heterogeneity issue. In practice, they use different data partitioning approaches to select samples similar to the target disks (i.e., disks to be predicted) and construct transfer learning models to predict the health status of the target disks. In the process of data partitioning, we reconsider the issue of minority samples and find that it is caused by unlimited data partitioning. To alleviate the impact of minority samples, we modify the naive partitioning criterion given by Pereira et al. (2017) to restrict data partitioning. Based on the above analysis, we determine the pattern for partitioning the data based on multiple issues, i.e., multi-property data partitioning (MDP). In terms of transfer learning, we propose to treat the partitioned data as the tasks for multi-task learning. As a result, we introduce model-agnostic meta-learning (MAML), which can provide unified learning for multiple tasks.

In this paper, we propose MDP-MAML to achieve disk failure prediction for multiple issues, and our contributions are described below:

1. To build a unified solution for disk failure prediction over multiple issues, we propose to treat the issues of heterogeneous disks, model aging, and environmental variation as the data heterogeneity problem, and apply the solution to heterogeneous disks to data heterogeneity.

2. We improve the solution to the issue to heterogeneous disks to solve multiple issues. On one hand, we propose a data partitioning method called MDP, which takes multiple issues into account and partitions data in a planned way. On the other hand, we treat each partitioned dataset as a task for learning, and create a unified prediction model to deal with multiple issues via a multi-task learning model.

3. We evaluate our proposed model on two public datasets and achieve state-of-the-art prediction performance with a low overhead for multiple issues.

2 Related works

2.1 Prediction models for issues

As the classic issues for disk failure prediction, disk heterogeneity (Rincón et al., 2017), model aging (Xiao et al., 2018), environmental variation (Lu et al., 2020), and minority samples (Zhang J et al., 2020b) have been studied in depth.

For the issue of heterogeneous disks, Rincón et al. (2017) selected the common SMART attributes for different disk models, which makes it possible to build a prediction model for multiple disk models. Sun et al. (2019) used an attribute distribution normalization algorithm to reduce the data distribution differences between different disk models. Botezatu et al. (2016) and Pereira et al. (2017) evaluated several transfer learning strategies and reduced the impact of the heterogeneous disks issue. Xie et al. (2018) proposed the optimized modeling engine (OME), employing the approach of cross-validation to determine the prediction model for a specific disk model. Zhang J et al. (2020a) proposed high-dimensional disk state embedding (HDDse) for generic failure detection, which employs a Siamese network framework trained on samples and makes predictions for multiple disk models.

For the issue of model aging, Jiang et al. (2019) updated their prediction model regularly to ensure the effectiveness of the model in long-term use. Xiao et al. (2018) designed an online labeling algorithm to annotate the healthy disks and employed Offline-RF (random forest) to guarantee prediction accuracy.

For the issue of environmental variation, Lu et al. (2020) introduced the location marker to disk samples, which helps the prediction models learn the differences between environments. In addition, they tested multiple classifiers and found that CNN-LSTM (convolutional neural network - long short-term memory) combined with a location marker (we refer to it as CNN-LSTM for convenience in this paper) performs close to the best in all situations. Luo et al. (2021) used the method of self-attention (Vaswani et al., 2017) to learn the correlation of disk states within the same environment, distinguishing the different environments implicitly.

For the issue of minority samples, Zhang J et al. (2020b) proposed transfer learning based failure prediction for minority disks (TLDFP), which combines samples from the majority disk model to build a

prediction model for the minority disk model, with the help of a transfer learning algorithm, i.e., Tradaboost. Zhou et al. (2022) proposed the active semi-supervised learning model ASLDP. ASLDP can make full use of unlabeled samples, and improve the generalization ability of the model under the condition of scarcity of labeled samples. In addition, HDDse and OME not only solve the issue of heterogeneous disks but also solve the issue of minority samples by combining samples from majority disk models.

2.2 Meta-learning

Contemporary ML models are typically trained from scratch for a specific task using a fixed learning algorithm designed by hand. These models achieve success in the area of disk failure prediction when disk samples are sufficient and homogeneous (Zhu et al., 2013; Rincón et al., 2017; Zhang J et al., 2020b; Shen et al., 2021). This excludes many applications where disk samples are rare or heterogeneous. Meta-learning has been widely used in the areas of multi-task learning (Buffelli and Vandin, 2022; Mao et al., 2022) and few-shot learning (Finn et al., 2017; Nichol et al., 2018; Frikha et al., 2021), and provides an alternative paradigm where an ML model gains experience over multiple tasks and uses this experience to improve its future learning performance (Hospedales et al., 2022). Most of the meta-learning models obtain over 0.96 accuracy on the Omniglot dataset (Lake et al., 2011) when classifying five unseen classes of images with only one sample for fine-tuning (Finn et al., 2017). Among all these models, MAML is well known for its generalizability and practicability (Finn et al., 2017). In this paper, we try to address the issues of failure prediction under the meta-learning paradigm, i.e., constructing disk samples into tasks for multi-task learning and using a few samples to fine-tune the prediction model.

3 Motivation

Although there are effective learning methods for the issues of heterogeneous disks, model aging, environmental variation, and minority samples, each method claims to be good at solving only one or two issues. We list them in Table 1. However, because issues do not arise in isolation, it is necessary to explore ways to solve multiple issues concurrently. As

Table 1 Applicability of the existing disk failure prediction methods

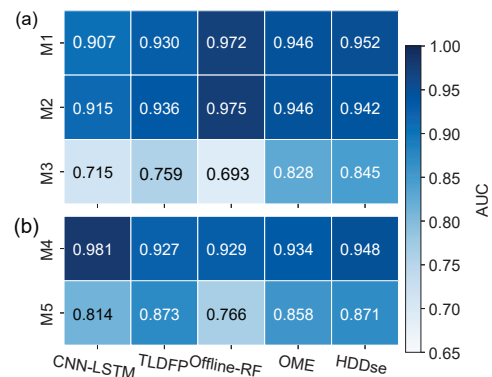
Method	Heterogeneous disks	Model aging	Environmental variation	Minority samples
CNN-LSTM (Lu et al., 2020)			✓	
TLDFP (Zhang J et al., 2020b)				✓
Offline-RF (Xiao et al., 2018)		✓		
OME (Xie et al., 2018)	✓			✓
HDDse (Zhang J et al., 2020a)	✓			✓

a result, we first simulate the scenario with multiple issues to evaluate the performance of the above methods. Note that this experiment is simply used to demonstrate and analyze the shortcomings of existing methods.

As shown in Table 2, we select three and two disk models from the Backblaze dataset (i.e., B) and open-source dataset (i.e., M) to build scenario 1 and scenario 2, respectively, where each property of the data selected from each dataset has been aligned. To simulate the scenario with the issues of heterogeneous disks, model aging, and minority samples, we further pick the disk data from the three models of dataset B over a five-year period and prepare a relatively small number of healthy and failed disks for M3. To simulate the scenario with the heterogeneous disks, environmental variation, and minority samples issues, we further pick the disk data from the two models of dataset M that are gathered from three server clusters and prepare a relatively small number of healthy and failed disks for M5. For each disk model, we randomly select 100 samples as the testing set and the rest as the training set.

As shown in Fig. 2, Offline-RF has the highest AUC value on M1, due to its optimization for the issue of model aging. Still, it fails to adapt to M3 and M5 because it relies on ample samples. TLDFP has the highest AUC on M5 due to its optimization of the minority samples issue. However, TLDFP does not obtain the claimed performance on M3. This is

because, for most minority disk models, it is difficult to find a suitable majority disk model for transfer learning (Zhang J et al., 2020a). CNN-LSTM takes the physical location into account, resulting in optimal results on M4. Nevertheless, its performance on other disk models is mediocre. Although OME and HDDse have relatively high AUC values on most disk models, they consume excessive computing resources, where HDDse increases the number of samples by a quadratic power, and OME needs to build at least three learning models for each disk model (see Section 6.5). In addition, they do not gain an overwhelming advantage on all disk models, where Offline-RF has the highest AUC on M1 and M2. These results demonstrate that no method can solve all issues well.

**Fig. 2** Prediction performance of different methods on two scenarios: (a) scenario 1; (b) scenario 2**Table 2** Selected disks for scenario simulation and performance evaluation

Scenario	Dataset	Period	Environment	Model	Number of disks	Number of failed disks
1	B	2016-01-01 to 2020-12-31	—	ST12000NM0007(M1)	38 739	1809
				ST4000DM000(M2)	36 158	3256
				ST6000DX000(M3)	1912	58
2	M	2016-05-18 to 2016-08-03	Six clusters	ST4000NM0033(M4)	54 665	583
				ST2000NM0033(M5)	6130	74

We execute evaluations on the Backblaze dataset and the dataset open-sourced by Lu et al. (2020), which are denoted as B and M , respectively

From a global perspective, M3 and M5 cannot be predicted well by all methods. On one hand, M3 has only 58 failed disks and spans five years. It demonstrates that without ample samples to update, even the state-of-the-art method cannot accurately predict. On the other hand, M5 represents the heterogeneous disks, environmental variation, and minority samples issues. This demonstrates that existing methods still have room for improvement in accuracy when dealing with multi-issue predictions.

Intuitively, we believe that solving the multi-issue prediction problem requires a unified learning framework. On one hand, we should explore uniform criteria for data preparation. On the other hand, we should seek a unified learning framework using the commonalities between issues. Nevertheless, existing methods have not been explored in depth in either of these areas.

4 Rationale

By analyzing the issues of heterogeneous disks, model aging, and environmental variation, we find that these issues can be attributed to the same problem, i.e., data heterogeneity, and can be learned using the same pattern, i.e., data partitioning and transfer learning. In addition, the issue of minority samples is special and always available because it is a result of unplanned data partitioning. Based on these reflections, we design a framework that includes a well-planned data partitioning method and a multi-domain transfer learning algorithm. This framework provides a unified solution for multiple issues, as discussed in Section 5.

4.1 Data heterogeneity

4.1.1 Identified issues

Heterogeneous disks: Conventionally, researchers argue that the issue of heterogeneous disks arises from differences in the distribution of model-specific SMART attributes (Botezatu et al., 2016). These distribution differences can introduce biases into the prediction model and reduce its prediction performance for each disk model.

Model aging: The previous research (Xiao et al., 2018) found that the distribution of SMART attributes can exhibit differences over time, even for the same disk model. This phenomenon results in

the model aging issue. The distribution differences also create biases in the prediction model and reduce its prediction performance in the future.

We can conclude that the properties of the period and disk models have similar effects on the distribution of SMART data. Therefore, we can boldly apply the period and disk models as equivalent properties to ML models because both properties can be used to distinguish between different SMART distributions. As a result, we argue that the issues of heterogeneous disks and model aging originate from the same problem, i.e., the difference in the distribution of SMART data, which we call data heterogeneity.

4.1.2 Emerging issue: environmental variation

The understanding of data heterogeneity can easily assimilate emerging issues. Recently, Lu et al. (2020) and Luo et al. (2021) found that location markers and the neighborhood disk information may affect the prediction accuracy of the ML models. This implies that disks in different environments, i.e., different physical locations and clusters, may manifest different distributions of SMART attributes.

We compare the SMART attributes of the disks from different environments, where we use the differences between server clusters to represent the environmental difference because the disks' physical location and workload differ between the server clusters. Fig. 3 shows the distribution differences of three SMART attributes between disks in the ay87a server cluster (E1) and ay84a server cluster (E2), although these disks come from the same disk model (i.e., ST4000NM0033) and the same period (i.e., 2016-05-18 to 2016-08-03). In addition, we use the samples from the disks in E1 and E2 to train four commonly used ML models and evaluate their prediction performance on disks from a new environment, i.e., the ay75c server cluster (E3). As shown in Fig. 4, the ML models perform worse on the disks from E3 than from other environments. This demonstrates that environmental variation creates differences in data distribution and results in low accuracy caused by training on samples with distribution biases, as we observe in the issue of heterogeneous disks. As a result, we believe that the environmental variation issue has the same learning process as the heterogeneous disks and model aging issues.

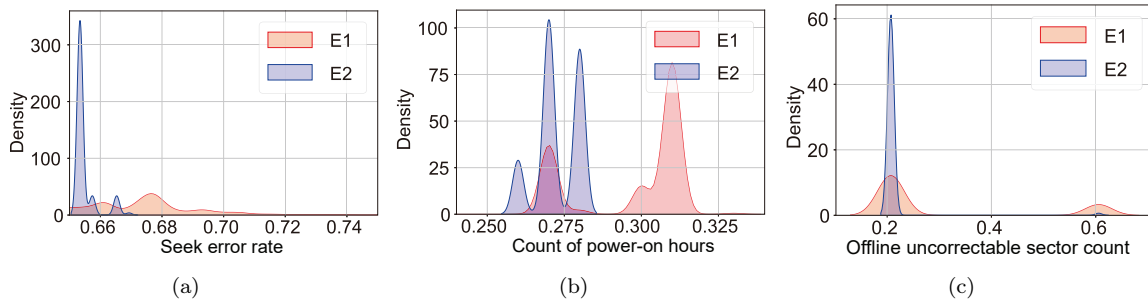


Fig. 3 Distribution of self-monitoring analysis and reporting technology (SMART) attributes in different environments: (a) SMART 7; (b) SMART 9; (c) SMART 198. E1 and E2 denote different environments

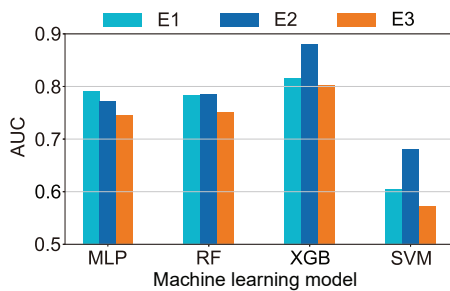


Fig. 4 Prediction performance of environments included or excluded in the training set. MLP: multi-layer perceptron; RF: random forest; XGB: XGBoost; SVM: support vector machine. E1–E3 denote different environments

4.2 Learning process for data heterogeneity

The issue of heterogeneous disks, as a special case of data heterogeneity, has been intensively studied (Botezatu et al., 2016; Pereira et al., 2017; Xie et al., 2018; Zhang J et al., 2020a). Through analysis of the commonality of these solutions, we find the process that learns the data to be similar to the target disks via data partitioning (i.e., preparation for training data) and the transfer learning model (Botezatu et al., 2016; Pereira et al., 2017). We believe that this learning process is equally valid for the issues of model aging and environmental variation.

Data partitioning: To learn the data from different disk models as independent samples, researchers partition the dataset into multiple subsets based on the disk model (Rincón et al., 2017; Xie et al., 2018). Each subset contains samples with the same disk model as the target disk, which results in a similar distribution of SMART attributes and good prediction performance of the ML model. This data partitioning approach can also be applied to model aging and environmental variation issues, such as partitioning data by month (Xiao et al., 2018) or by the

environment. However, data partitioning methods are not readily available for multi-issue scenarios. In Section 5.1, we propose MDP to address the challenge of data partitioning in such scenarios.

Transfer learning: Transfer learning uses experience from a source task to improve learning (speed, data efficiency, and accuracy) on a target task (Hospedales et al., 2022). When applied to the disk heterogeneity issue, representative samples similar to the target model disk samples are selected as transfer sources, and the classifier is used for learning these samples (Zhang J et al., 2020b). However, the applicability of this method is limited by the transfer learning methods used for prediction models, which can handle only single-domain tasks. When multiple issues need to be addressed, such as heterogeneous disks and environmental variation, the single-domain model can result in sub-optimal prediction results. Specifically, disks from the same model may come from different environments, and vice versa. However, the single-domain model can select either similar environments or similar disk models as its source domain. To address this problem, in Section 5.2, we introduce a multi-domain transfer learning model that can learn from multiple domains, including the disk model, environment, and time period.

4.3 Minority samples

In the data center, the number of some model disks is dramatically smaller than that of others, and we call these disks “minority disks.” Using traditional ML algorithms with the training data of minority disks would dramatically increase the risk of over-fitting or poor generalization, which would weaken the performance of predictive models and seriously affect the reliability of the storage

system (Zhang J et al., 2020b). Furthermore, we will give a more specific explanation in this subsection. First, the issue of minority samples does not exist independently; it is caused by unplanned data partitioning and must therefore be considered in the context of the multi-issue scenario. Second, we draw a qualitative conclusion; i.e., it is the number of failed disks rather than the number of disks that determines the minority sample status.

In practice, public datasets provide a large number of disks. However, data partitioning will narrow the dataset to avoid the data heterogeneity issue. Consequently, it may result in fewer available samples and lower prediction accuracy. Although the transfer learning method has been proven effective in handling the issue of minority samples, it is still tricky in appropriately partitioning the data for the balance of the data heterogeneity and minority samples issues.

Intuitively, to perform data partitioning, we need to ensure that each partitioned subset should contain the right number of disks for learning accuracy. Zhang J et al. (2020b) found that the issue of minority samples arises when the number of disks in the subset is smaller than 1500. Specifically, 1500 is the minimum number of disks that causes the generalization error (i.e., difference between testing loss and training loss) to converge. However, as shown in Fig. 5, the convergence points of the generalization error curves for different disk models differ significantly in the scenario with heterogeneous disks, which results in the previous approach being no longer applicable. Xie et al. (2018) argued that the number of failed disks is crucial to the convergence of the ML model. Inspired by this insight, we redraw the generalization error curves using the number of failed disks as the horizontal coordinate. The results are shown in Fig. 6. It is clear that the generalization error for each disk model converges around 150 failed disks. Combining this with the conclusion from Xie et al. (2018), we determine that the minority samples problem refers to having a too small number of failed disk samples of certain disk models. As a result, to balance minority samples and data heterogeneity issues, we design the data partitioning method MDP to ensure that the partitioned subset should contain a minimum number of failed disks rather than the number of disks. The minimum number of failed disks that should be con-

tained within a partitioned subset is described in Section 6.3.

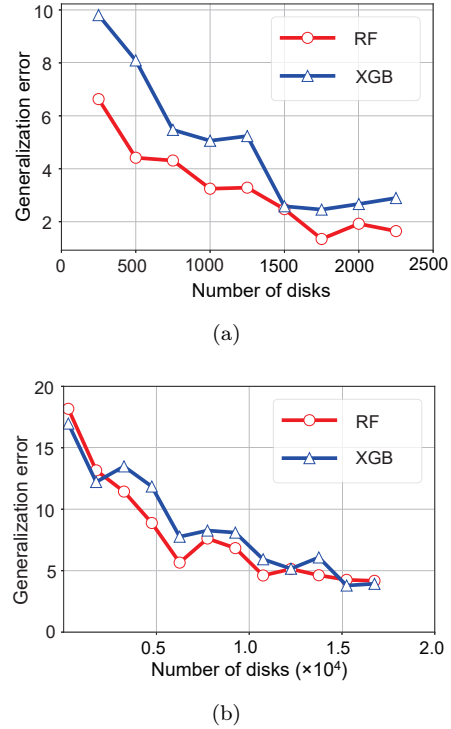


Fig. 5 Generalization error vs. the number of disks: (a) ST4000DM000; (b) ST12000NM0008

5 Design of MDP-MAML

To implement prediction for the issue of data heterogeneity using minority samples, we design the model in terms of data partitioning and multi-domain learning and propose MDP-MAML. MDP-MAML consists of a multi-property data partitioning method MDP and a multi-task learning model MAML (Fig. 7). The white arrows represent the data preparation process, and the blue arrows represent the process of MAML learning and prediction for multiple tasks. During the data preparation process, MDP partitions the mixture of samples into multiple subsets using multiple properties (e.g., disk model, period, and environment), which reduces the data heterogeneity within each subset. During the learning and prediction processes, MAML treats the partitioned subsets (e.g., M1E1T1 and M1E2T1) as tasks for multi-task learning, extracts representative features from all the tasks, and uses them to make predictions on the new disk model, new environment, and latest period (e.g., M4E3Test). In addition, we

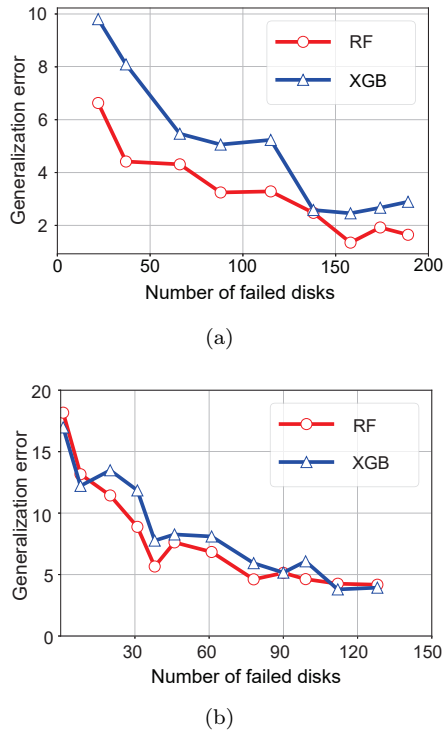


Fig. 6 Generalization error vs. the number of failed disks: (a) ST4000DM000; (b) ST12000NM0008

parallelize the multi-task learning process to reduce the training overhead.

5.1 MDP

There are two challenges to data partitioning in multi-issue scenarios. The first challenge is uneven

partitioning. Zhang J et al. (2020b) verified that performing data partitioning by disk models on the Backblaze dataset will result in a large gap between the partitioned subsets. As shown in Table 3, the subset of M6 contains 2993 failed disks, whereas the subset of M1 contains only 24 failed disks. In addition, the period of M6 ranges from 2016 to 2020, which means that M6 may suffer from the issue of model aging, whereas M1 may suffer from the issue of minority samples. If we further divide M1 by time, then the subset of failed disks within the partitioned subset will further shrink, leading to a more severe minority samples issue. The second challenge is partitioning for the property with continuous values. Specifically, there is no clear partition granularity for the time period property. The sample can be partitioned by week, month, or year. However, even within the same period, the number of failed disks varies dramatically from one model to another. Partitioning the data without a plan will only worsen the problem of uneven partitioning.

As a result, we propose a data partitioning method named MDP that partitions data using multiple properties. MDP continuously selects larger subsets for partitioning based on multiple properties until the number of failed disks in each subset is smaller than a critical value. As shown in Fig. 8, the large cuboid at the top left represents a sample collection with three dimensions: disk model, environment, and period. Each small cube within the cuboid represents one disk sample. When preparing

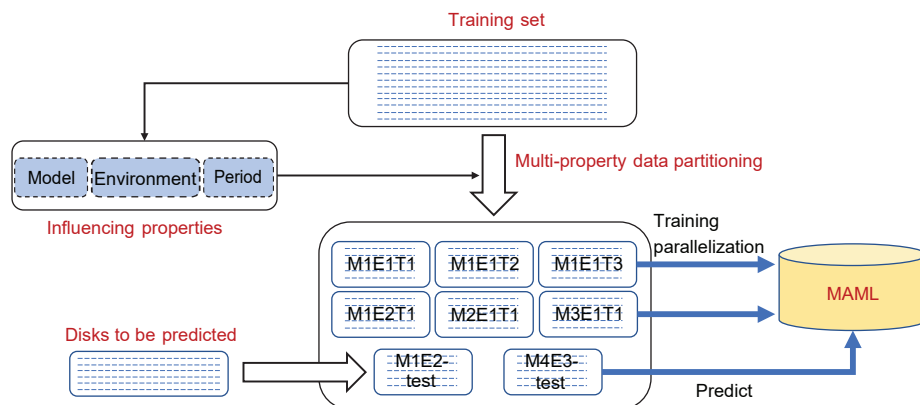


Fig. 7 Overview of the model. By multi-property data partitioning (MDP), the training set is partitioned into multiple subsets by the properties corresponding to the issues. For example, M1E1T1 denotes the subset consisting of disk samples from a certain disk model, a certain environment, and a certain period. M denotes the disk model, E denotes the environment, and T denotes the time period. In addition, because the disks to be predicted are in a later period than the training set, we use “test” to denote the period of these disks. References to color refer to the online version of this figure

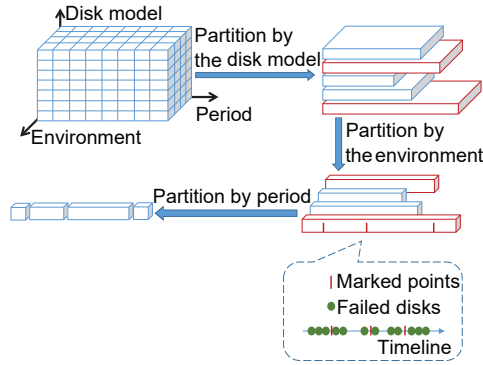


Fig. 8 Multi-property data partitioning (MDP). We partition data by disk models, environments, and periods, separately. The red cuboids represent the subsets that need to be further partitioned and the blue cuboids represent the subsets for multi-task learning. When partitioning by periods, we determine the partitioning according to the number of failed disks in the partitioned subset. References to color refer to the online version of this figure

the training set, we first partition the mixed samples into multiple subsets by disk models and select the subsets (red cuboids) that contain more than 70 failed disk samples (this number is determined in Section 6.3). Then, we continue to partition the selected subsets by environments and select the subsets that contain more than 70 failed disk samples. Finally, we partition the selected subsets by period. To this end, we construct a timeline according to the recording time of the disk samples and mark every 70 failed disk samples. We then partition the samples according to the marked points on the timeline. In this way, disk samples from different models and environments are partitioned by dynamic period intervals to ensure even partitioning. The remaining blue cuboids after each partition are the subsets for multi-task learning. Note that the partitioned subsets consist of both healthy and failed disk samples, with the former accounting for the majority of the population. To relieve the data imbalance issue, we leverage the commonly used under-sampling method (He and Garcia, 2009) to under-sample the healthy disk samples, resulting in different ratios of failed to healthy samples ranging from 1:1 to 1:50. In the final training set, we set this ratio to 1:3, which leads to superior prediction accuracy. Note that using only under-sampling or over-sampling cannot effectively address the issue of minority samples, and it is a common practice to leverage transfer learning algorithms with the help of data from other types of

disks. This approach refers to the discussion in Section 4.3, i.e., using the number of failed disks as the metric to quantify data partitioning and balance the minority samples and data heterogeneity issues. It also guarantees that large subsets will be partitioned based on multiple properties, whereas small subsets will retain a certain number of failed disks.

5.2 Multi-task learning algorithm—MAML

To tackle the issue that single-domain transfer learning cannot deal with the issues associated with multiple properties, we resort to the multi-domain learning method. Multi-domain learning is a kind of multi-task learning that is trained on multiple domains to learn models for each target domain (Wang et al., 2021). As a result, we propose to leverage all the partitioned subsets and treat each partitioned subset as a task for multi-task learning.

MAML (Finn et al., 2017) is a meta-learning algorithm commonly used in the field of multi-task learning. MAML optimizes the process of “fine-tuning on a task” on multiple tasks to achieve fast adaptation to unknown tasks. When learning on a task, MAML is concerned with the process of fine-tuning on that task, instead of the domain of that task. This allows MAML to extract the common features of tasks from multiple domains to improve generalization capabilities (Hospedales et al., 2022). In addition, because the fine-tuning process is continuously optimized during training, the number of samples required to perform fine-tuning is continuously reduced, which further alleviates the minority samples issue. In this study, we consider each partitioned subset as a task (i.e., prediction on a model, an environment, and a period) and target disks coming from a new disk model, a new environment, and a future period as the unknown task.

5.2.1 Training process

MAML optimizes the process of “fine-tuning on a task” on multiple tasks to achieve fast adaptation to unknown tasks. Formally, the process can be written as

$$f_{\theta'} = f_{\theta} - \eta \nabla f_{\theta}(S), \quad (1)$$

where S represents data for fine-tuning and η represents the learning rate of fine-tuning. Then, MAML calculates the loss of this process on each task and chooses a gradient descent algorithm to optimize this

loss. Formally, it can be written as

$$f_\phi = f_\phi - \alpha \sum_{i=1}^K \nabla f_{\theta'_i}(Q_i), \quad (2)$$

where Q_i represents data for evaluation, α represents the learning rate of gradient descent, and K represents the number of training tasks. Note that f_ϕ and f_θ are learners who share the same architecture. The complete training process is shown in Algorithm 1.

Algorithm 1 Training process of MAML

```

1: Initialize task set  $\{\text{task}_i\}$  and meta-learner  $f_\phi$ 
2: Initialize learning rates  $\alpha$  and  $\eta$ 
3: while not done do
4:   for all taski do
5:      $f_{\theta_i} = f_\phi$ 
6:      $f_{\theta'_i} = f_{\theta_i} - \eta \nabla f_{\theta_i}(S_i)$ 
7:     Calculate loss of fine-tuning  $f_{\theta'_i}(Q_i)$ 
8:   end for
9:    $f_\phi = f_\phi - \alpha \sum_{i=1}^K \nabla f_{\theta'_i}(Q_i)$ 
10: end while

```

5.2.2 Prediction process

The MAML prediction process is the process of fine-tuning. Specifically, data used for fine-tuning, i.e., S_{test} , are newly labeled or selected from the training set. It is the same as the samples to be predicted, i.e., Q_{test} , in terms of the disk model, environment, and period. Formally, we have

$$f_{\phi^{*'}} = f_{\phi^*} - \eta \nabla f_{\phi^*}(S_{\text{test}}), \quad (3)$$

where f_{ϕ^*} represents the well trained meta-learner. Then, $f_{\phi^{*'}}$ predicts the status of the disk samples in Q_{test} and reports the prediction result, i.e.,

$$\bar{y} = f_{\phi^{*'}}(Q_{\text{test}}). \quad (4)$$

5.2.3 Parallelization implementation

For most ML models optimized by gradient descent, each round of gradient descent depends on the result of the previous gradient descent, making parallelization difficult to implement. However, according to Algorithm 1, the input (S_i) and output ($f_{\theta'_i}(Q_i)$) of each f_{θ_i} are independent, so there is a possibility for parallelization implementation. As a result, we can clone multiple f_{θ_i} to simultaneously compute the expectation loss $f_{\theta'_i}(Q_i)$ for each task. Note that the number of cloned f_{θ_i} is limited by the GPU's memory.

6 Evaluation

In this section, we compare MDP-MAML's prediction performance and overhead with those of several state-of-the-art prediction models on two datasets. In addition, we evaluate the effectiveness of MDP using ablation experiments. The experiments were executed on a Linux server with a 20-core 2.4 GHz CPU, 128 GB RAM, and NVIDIA GPU 3090.

6.1 Dataset and data preprocessing

Dataset: We built datasets based on samples collected from the Backblaze dataset and the dataset open-sourced by Lu et al. (2020); both datasets consist of SMART attributes collected from real-world data centers. We denote the above two datasets as B and M , respectively. The information on the built datasets is shown in Table 3. For dataset B , we used the samples from six disk models with a period spanning four years as the training set, testing the prediction accuracy for three new disk models that emerged in the next year. For dataset M , we used the samples from three disk models in four environments as the training set to test the performance of the prediction models for three new disk models in three new environments.

Feature selection: For each disk in our datasets, 30 SMART attributes were reported. Each attribute contains two values of interest, including a raw value and a normalized value. However, as some SMART attributes are irrelevant to disk failures, we needed to refine them. The SMART attributes we selected (Table 4) were the same as used by Zhang J et al. (2020b), to evaluate the performance of the prediction model on multiple disk models.

Data normalization: Because different SMART attributes have diverse value intervals, we applied data normalization to ensure a fair comparison. The normalization function (Xiao et al., 2018) used in our approach is

$$x' = \frac{x - x_{\min}}{x_{\max} - x_{\min}}, \quad (5)$$

where x is the original value of a feature, and $x_{\max}$ and $x_{\min}$ are the maximum and minimum values of the feature in our datasets, respectively.

Sample labeling: Zhang J et al. (2020a) observed that most attributes undergo significant changes around seven days before failures. Thus,

Table 3 Overview of the datasets

Dataset	Period	Environment	Model	Number of disks	Number of failed disks
B_{train}	2016-01-01 to 2019-12-31	–	ST1000NM0086(M1)	1256	24
			ST6000DX000(M2)	1725	55
			ST8000DM002(M3)	10 211	362
			ST8000NM0055(M4)	15 038	447
			ST12000NM0007(M5)	26 762	1482
			ST4000DM000(M6)	31 184	2993
B_{test}	2020-01-01 to 2020-12-31	–	ST12000NM001G(M7)	7161	34
			ST12000NM0008(M8)	19 527	146
			HGST HMS5C4040BLE640(M9)	12 744	35
M_{train}	2016-05-18 to 2016-08-03	E1, E2, E3, E4	ST600MM0006(M10)	5092	47
			ST4000NM0033(M11)	44 266	458
			ST32000645NS(M12)	48 928	469
M_{test}	2016-05-18 to 2016-08-03	E5	ST2000NM0011(M13)	4215	30
		E6	ST2000NM0033(M14)	5713	62
		E7	HGST HUS724020ALA640(M15)	4768	44

B and M denote the Backblaze dataset and the open-source dataset, respectively. M1–M15 denote different disk models and E1–E7 denote different server clusters

for a failed disk, continuous samples in the seven-day period before the actual failure were labeled as failed ($y = 1$), and other samples were labeled as healthy ($y = 0$). For healthy disks, all samples were labeled as healthy.

Table 4 SMART attributes selected for evaluation

Attribute ID	Attribute name	Attribute type
1	Raw read error rate	Normalized & raw
3	Spin-up time	Normalized
5	Reallocated sector count	Normalized & raw
7	Seek error rate	Normalized & raw
9	Power-on hours	Normalized & raw
184	I/O error detection and correction	Normalized & raw
187	Reported uncorrectable errors	Normalized & raw
188	Command timeout	Raw
189	High fly writes	Normalized & raw
190	Airflow temperature	Normalized & raw
193	Load/unload cycle count	Normalized & raw
194	Temperature	Normalized & raw
197	Current pending sector count	Normalized & raw
198	Offline uncorrectable sector count	Normalized & raw
240	Head flying hours	Raw
241	Total LBAs written	Raw
242	Total LBAs read	Raw

LBA: logical block addressing

6.2 Metrics

We adopted three metrics which are commonly used for evaluating the capability of a prediction model to report the results in our experiments:

False detection rate (FDR), also called the recall rate, captures the proportion of failed disks that are correctly predicted as failed. We measured the FDRs under the constraint that the false alarm rates (FARs) were around 1.0%.

False alarm rate (FAR) represents the proportion of healthy disks that are falsely predicted as failed. We measured the FARs under the constraint that the FDRs were around 0.95.

AUC represents the area under the receiver operating characteristic curve (Ling et al., 2003). It is a performance measurement for classification problems at various thresholds. In disk failure prediction, a high AUC value means that the model does well in distinguishing between failed and healthy disks.

6.3 Hyper-parameter configuration

The hyper-parameters involved in MDP-MAML are the architecture parameters of learner f_ϕ , the maximum number of training epochs, the number of gradient descents for fine-tuning, and the lower bound of the number of failed disks for performing data partitioning. For the architecture, we used a neural network with four hidden layers with sizes

of 256, 128, 128, and 64, each including batch normalization and ReLU activation, followed by a linear layer and a sigmoid function (Finn et al., 2017). The meta-learner and fine-tuning learning rates were 0.001 and 0.0005, respectively, and the maximum number of training epochs was 1000. We used five steps of gradient descent for fine-tuning. In terms of the partitioning lower bound, as shown in Fig. 9, according to the best AUC value in the curve depicted by different lower bounds, we selected the appropriate value of the lower bound, i.e., 70. In practice, we stopped data partitioning when the number of failed disks in a subset was <70 . Note that we used AUC as the metric to determine the best lower bound for partitioning to maximize the accuracy of MDP-MAML. Fig. 6 shows the generalization error, which is a metric used for qualitative analysis to infer the indicator of minority samples states, i.e., the number of failed disks. In addition, the value of 70 obtained here is not universally applicable to all algorithms. Theoretically, the stronger the algorithm's generalization capability, the smaller this value should be.

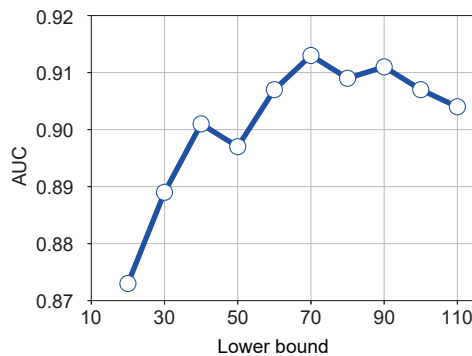


Fig. 9 AUC vs. the partitioning lower bound

6.4 Comparison with state-of-the-art methods

We prepared datasets containing the issues of heterogeneous disks and environmental variation, i.e., M_{train} and M_{test} , and the issues of heterogeneous disks and model aging, i.e., B_{train} and B_{test} . On these datasets, we compared MDP-MAML with HDDse (Zhang J et al., 2020a), OME (Xie et al., 2018), TLDFP (Zhang J et al., 2020b), Offline-RF (Xiao et al., 2018), and CNN-LSTM (Lu et al., 2020) in terms of applicability. For the parameter setting, we deployed HDDse and CNN-LSTM according to their original papers, respectively, while deploying TLDFP, Offline-RF, and OME according to the pa-

rameters acquired by grid search. Because TLDFP and Offline-RF do not discuss their applicability to new disk models or new environments, we deployed these methods using 10 disks with labels in the testing set to ensure performance. For fairness, we added these 10 disks to the training set to train HDDse, OME, and CNN-LSTM. For MDP-MAML, we used these 10 disks for fine-tuning, i.e., considering these labeled disks as S_{test} .

Table 5 lists the average prediction performance of the above methods. Compared to the state-of-the-art methods, MDP-MAML can improve AUC and FDR from 0.85 to 0.89 and from 0.85 to 0.91, respectively, while reducing FAR from 4.88% to 2.85%. Although MDP-MAML improved FDR, as we will see in the next subsections, it outperformed other methods in all scenarios, proving its generalizability. In the following subsections, we will show the prediction performance from each issue's perspective.

Table 5 Average performance of the prediction model on the testing set

Model	AUC	FDR	FAR (%)
MDP-MAML	0.89	0.91	2.85
HDDse	0.83	0.84	5.65
OME	0.85	0.85	4.88
TLDFP	0.81	0.79	6.48
CNN-LSTM	0.77	0.74	8.68
Offline-RF	0.71	0.62	12.15

6.4.1 Heterogeneous disks

For the issue of heterogeneous disks, we evaluated each method's prediction performance via six new disk models, i.e., M7, M8, and M9 in dataset B and M13, M14, and M15 in dataset M . As shown in Fig. 10a, MDP-MAML outperformed HDDse and OME in dealing with the issue of heterogeneous disks and obtained the highest AUC value on the two datasets. The same results can be seen in Figs. 10b and 10c. We attribute the advantage of MDP-MAML to its applicability to multiple issues. For example, there was also a model aging issue in the testing set of dataset B because the periods of samples spanned four years. Similarly, the disks in dataset M came from multiple environments, implying the existence of the environmental variation issue. Note that compared to working on Seagate disks, MDP-MAML obtained lower accuracy on M9 and M13 from Hitachi, because we did not use Hitachi disks

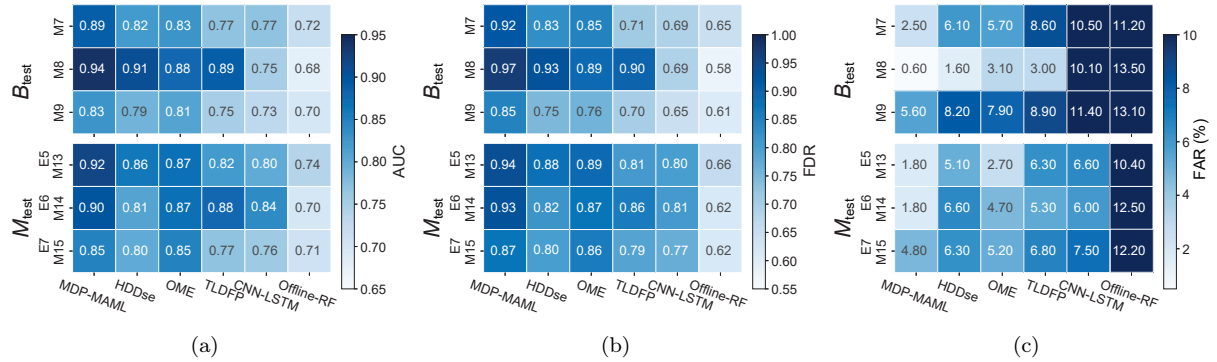


Fig. 10 Comparisons of performance with state-of-the-art disk failure prediction methods on the issues of heterogeneous disks and environmental variation: (a) AUC; (b) FDR; (c) FAR

for training. Nevertheless, our prediction results were better than those of others. In addition, multiple issues inhibited the strength of the traditional methods in their respective expertise. For example, Offline-RF did not obtain the results reported in its original paper. We believe that Offline-RF needs to build a random forest for each disk model. However, because only 10 disks with labels can be used for training, Offline-RF suffered from the issue of minority samples.

6.4.2 Environmental variation

For the issue of environmental variation, we mainly observed the prediction performance for three new environments, i.e., E5, E6, and E7. As shown in Fig. 10, MDP-MAML had the best prediction performance compared to state-of-the-art methods on E5, E6, and E7. In addition, CNN-LSTM did not have the expected results because it suffered from the heterogeneous disks issue.

6.4.3 Model aging

For the issue of model aging, we mainly observed the results for M7 and M8 whose periods spanned from 2020-01-01 to 2020-12-31. MDP-MAML and HDDse used B_{train} to build prediction models. Offline-RF and HDDse used the samples from the previous month to update models, which began with the second month. MDP-MAML uses the samples from the previous month as S_{test} for fine-tuning and the samples from the current month as Q_{test} for evaluation. As shown in Figs. 11a and 11b, MDP-MAML gained the highest average performance and the most stable prediction performance

on M7 and M8. We attribute this to insufficient samples for model updating on M7. In contrast, because MAML optimizes this process of fine-tuning, it can perform fine-tuning with a small number of samples when predicting the samples from a new period.

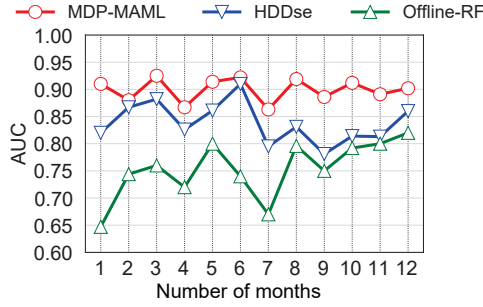
6.5 Comparisons in overhead

To demonstrate that MDP-MAML does not incur significant overhead while yielding better performance, we compared the time overhead of MDP-MAML with those of OME and HDDse in three phases, i.e., data preparation, model training, and prediction. We performed data preparation and model training on B_{train} , which contained samples of 86 176 disks. By comparing HDDse and MDP-MAML, we calculated the average time for predicting 1000 samples.

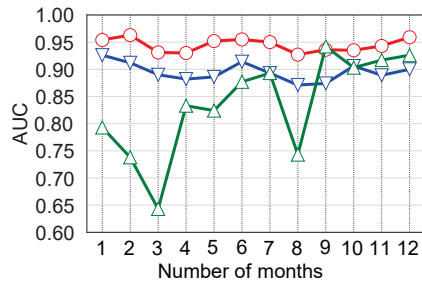
Data preparation: We enumerated the operations required by each method's training data preparation. As shown in Table 6, SMART attribute selection and data normalization were required by all the methods. HDDse and OME partitioned data by using disk models, while MDP-MAML partitioned data using multiple properties. HDDse required an additional operation to construct sample pairs.

Table 6 Operations required by the data preparation of each method

Operation	OME	HDDse	MDP-MAML
Attribute selection	✓	✓	✓
Partition by disk models	✓	✓	
Partition by multiple properties			✓
Data normalization	✓	✓	✓
Constructing sample pairs		✓	



(a)



(b)

Fig. 11 AUC for the issue of model aging on M7 (a) and M8 (b)

Model training and prediction: During the model training phase, OME needed to train three ML models for each disk model, while MDP-MAML and HDDse trained one ML model shared by all the disks. During the testing phase, OME required additional cross-validation operations to find the best-fit ML model for each disk model. In addition, HDDse needed to conduct an additional output comparison to determine the health status of each disk.

As shown in Fig. 12, in the data preparation phase, HDDse consumed the most time due to the additional sample-pairing operation, whose time complexity was $O(n^2)$, where n represents the number of samples in the training set. Although MDP-MAML partitioned data using multiple attributes, its time complexity was $O(n)$, the same as that of partitioning data only by disk models. In the training phase, MDP-MAML can accelerate the training by parallelization. In contrast, HDDse had a long training time due to its large number of sample pairs, while OME needed to build three ML models for each disk model. In the testing phase, OME's cross-validation and HDDse's status decision incurred additional overhead. As a result, MDP-MAML can address multiple issues while maintaining a low overhead.

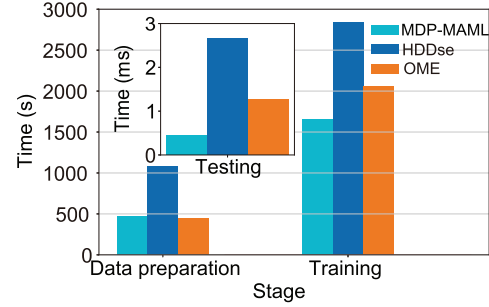


Fig. 12 Comparison in overhead

6.6 Effectiveness of MDP

To illustrate the effectiveness of MDP, we showed the results of AUC using different data partitioning strategies, i.e., implementing partitioning based on different combinations of properties. We selected M_{train} and M_{test} to conduct this experiment. Because the samples in dataset M spanned fewer than three months, we evaluated performance every half month to simulate model aging issues. This differs from the commonly used one-month evaluation interval. As shown in Table 7, the best AUC result was achieved when we executed MDP for three properties. We believe that fine-grained partitioning can further improve the quality of prediction models. Furthermore, MAML can support multi-task learning after fine-grained partitioning. In addition, MDP is effective for data partitioning in dealing with each issue.

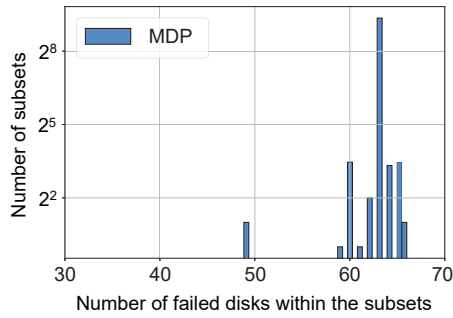
Note that we plot the distribution of the number of failed disks contained in the subsets to show the evenness of the data. Compared to the traditional partition method, which considers merely disk models, MDP can acquire more evenly distributed data. As shown in Fig. 13, most of the subsets partitioned by MDP had 60–70 failed disks, whereas the number of counterparts in the subset partitioned by the traditional partition method varied greatly. As a result, MDP can deal with multiple properties and ensure that the partitioned subsets have sufficient failed disks, addressing the challenge of data partitioning in the multi-issue scenario.

6.7 Lead time of MDP-MAML

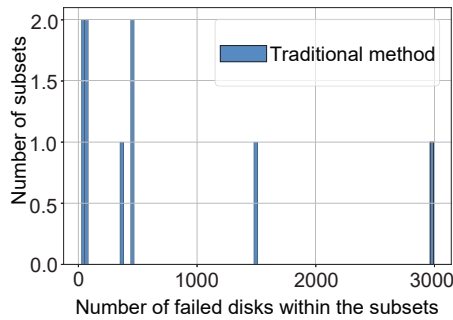
In addition to the accuracy of our prediction, it is crucial to consider the amount of time that users are given to back up their data. Another aspect worth examining is how much lead time

Table 7 Effectiveness of MDP

Property			AUC (mean)
Model	Period	Environment	
✓	✓	✓	0.912
	✓	✓	0.793
✓		✓	0.907
✓	✓		0.856



(a)



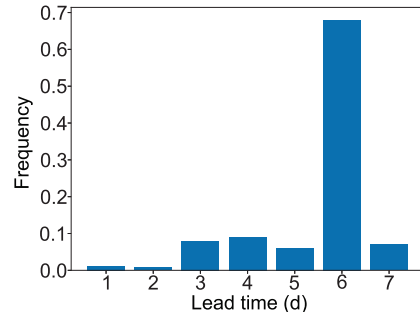
(b)

Fig. 13 Distribution of the number of failed disks within the partitioned subsets: (a) MDP; (b) traditional method

MDP-MAML can provide for detecting an impending failure. Fig. 14 displays the distribution of lead time for correct predictions. It is clear that MDP-MAML can predict potentially failed disks about six days in advance. We attribute this to the fact that we labeled the samples of the last seven days before the disk failure occurred as failures during training. We believe that the lead time of MDP-MAML can be improved by changing the principle of failed sample labeling, but this may come at the expense of lower accuracy.

7 Conclusions and future work

In this paper, we propose a failure prediction model, i.e., MDP-MAML, to handle multiple issues. Our main contributions include the following: (1) We

**Fig. 14 Lead time of MDP-MAML**

discover the common nature between the issue of heterogeneous disks, model aging, and environmental variation, i.e., data heterogeneity. (2) We summarize the commonality of different solutions to data heterogeneity and build a unified solution pattern for data heterogeneity, i.e., partitioning data and transfer learning. (3) We propose MDP and introduce MAML to overcome the shortcomings of data partitioning and transfer learning methods used by the disk failure prediction solutions in the multi-issue scenario. Our experiments on datasets gathered from real-world data centers demonstrate that MDP-MAML outperforms its state-of-the-art counterparts in terms of prediction performance and overhead.

With the growing use of solid-state drives (SSDs) in data centers, prediction of their failure has become a hot issue. Studies have shown that SSD failures rely not only on SMART attributes, but also on other disk- and system-level logs (Zhang YQ et al., 2023). To ensure the applicability of MDP-MAML to SSDs, we need to use information beyond SMART attributes to enhance its generalizability.

Contributors

Yunchuan GUAN designed the research and conducted the experiments. Yunchuan GUAN and Yu LIU drafted the paper. Ke ZHOU helped organize the paper. Qiang LI, Tuanjie WANG, and Hui LI provided the data and funded the research. Yunchuan GUAN revised and finalized the paper.

Compliance with ethics guidelines

Yunchuan GUAN, Yu LIU, Ke ZHOU, Qiang LI, Tuanjie WANG, and Hui LI declare that they have no conflict of interest.

Data availability

The data that support the findings of this study are

available from the corresponding author upon reasonable request.

References

- Backblaze, 2018. The Backblaze Hard Drive Data and Stats. <https://www.backblaze.com/b2/hard-drive-test-data.html> [Accessed on Oct. 18, 2022].
- Botezatu MM, Giurgiu I, Bogojeska J, et al., 2016. Predicting disk replacement towards reliable data centers. Proc 22nd ACM SIGKDD Int Conf on Knowledge Discovery and Data Mining, p.39-48. <https://doi.org/10.1145/2939672.2939699>
- Buffelli D, Vandin F, 2022. Graph representation learning for multi-task settings: a meta-learning approach. Int Joint Conf on Neural Networks, p.1-8. <https://doi.org/10.1109/IJCNN55064.2022.9892010>
- Finn C, Abbeel P, Levine S, 2017. Model-agnostic meta-learning for fast adaptation of deep networks. Proc 34th Int Conf on Machine Learning, p.1126-1135.
- Frikha A, Krompaß D, Köpken HG, et al., 2021. Few-shot one-class classification via meta-learning. Proc AAAI Conf on Artificial Intelligence, p.7448-7456. <https://doi.org/10.1609/aaai.v35i8.16913>
- Han SJ, Lee PPC, Xu F, et al., 2021. An in-depth study of correlated failures in production SSD-based data centers. Proc 19th USENIX Conf on File and Storage Technologies, p.417-429.
- He HB, Garcia EA, 2009. Learning from imbalanced data. *IEEE Trans Knowl Data Eng*, 21(9):1263-1284. <https://doi.org/10.1109/TKDE.2008.239>
- Hospedales T, Antoniou A, Micaelli P, et al., 2022. Meta-learning in neural networks: a survey. *IEEE Trans Patt Anal Mach Intell*, 44(9):5149-5169. <https://doi.org/10.1109/TPAMI.2021.3079209>
- Jiang TM, Zeng JF, Zhou K, et al., 2019. Lifelong disk failure prediction via GAN-based anomaly detection. *IEEE 37th Int Conf on Computer Design*, p.199-207. <https://doi.org/10.1109/ICCD46524.2019.00033>
- Lake BM, Salakhutdinov R, Gross J, et al., 2011. One shot learning of simple visual concepts. Proc 33rd Annual Meeting of the Cognitive Science Society, p.2568-2573.
- Ling CX, Huang J, Zhang H, et al., 2003. AUC: a statistically consistent and more discriminating measure than accuracy. Proc 18th Int Joint Conf on Artificial Intelligence, p.519-524.
- Lu SD, Luo B, Patel T, et al., 2020. Making disk failure predictions smarter! Proc 18th USENIX Conf on File and Storage Technologies, p.151-167.
- Luo C, Zhao P, Qiao B, et al., 2021. NTAM: neighborhood-temporal attention model for disk failure prediction in cloud platforms. Proc Web Conf, p.1181-1191. <https://doi.org/10.1145/3442381.3449867>
- Mao YR, Wang ZK, Liu WW, et al., 2022. MetaWeighting: learning to weight tasks in multi-task learning. Findings of the Association for Computational Linguistics: ACL 2022, p.3436-3448. <https://doi.org/10.18653/v1/2022.findings-acl.271>
- Nichol A, Achiam J, Schulman J, 2018. On first-order meta-learning algorithms. <https://arxiv.org/abs/1803.02999>
- Pereira FLF, dos Santos Lima FD, de Moura Leite LG, et al., 2017. Transfer learning for Bayesian networks with application on hard disk drives failure prediction. Brazilian Conf on Intelligent Systems, p.228-233. <https://doi.org/10.1109/BRACIS.2017.64>
- Pinheiro E, Weber WD, Barroso LA, 2007. Failure trends in a large disk drive population. Proc 5th USENIX Conf on File and Storage Technologies, p.17-28.
- Rincón CA, Pâris JF, Vilalta R, et al., 2017. Disk failure prediction in heterogeneous environments. Int Symp on Performance Evaluation of Computer and Telecommunication Systems, p.1-7. <https://doi.org/10.23919/SPECTS.2017.8046776>
- Shen J, Ren YJ, Wan J, et al., 2021. Hard disk drive failure prediction for mobile edge computing based on an LSTM recurrent neural network. *Mob Inform Syst*, 2021:1-12. <https://doi.org/10.1155/2021/8878364>
- Sun XY, Chakrabarty K, Huang RR, et al., 2019. System-level hardware failure prediction using deep learning. Proc 56th ACM/IEEE Design Automation Conf, p.1-6. <https://doi.org/10.1145/3316781.3317918>
- Vaswani A, Shazeer N, Parmar N, et al., 2017. Attention is all you need. Proc 31st Int Conf on Neural Information Processing Systems, p.6000-6010.
- Wang JD, Lan CL, Liu C, et al., 2021. Generalizing to unseen domains: a survey on domain generalization. Proc 30th Int Joint Conf on Artificial Intelligence, p.4627-4635. <https://doi.org/10.24963/ijcai.2021/628>
- Xiao J, Xiong Z, Wu S, et al., 2018. Disk failure prediction in data centers via online learning. Proc 47th Int Conf on Parallel Processing, p.1-10. <https://doi.org/10.1145/3225058.3225106>
- Xie YW, Feng D, Wang F, et al., 2018. OME: an optimized modeling engine for disk failure prediction in heterogeneous datacenter. *IEEE 36th Int Conf on Computer Design*, p.561-564. <https://doi.org/10.1109/ICCD.2018.00089>
- Zhang J, Huang P, Zhou K, et al., 2020a. HDDse: enabling high-dimensional disk state embedding for generic failure detection system of heterogeneous disks in large data centers. Proc USENIX Annual Technical Conf, p.111-126.
- Zhang J, Zhou K, Huang P, et al., 2020b. Minority disk failure prediction based on transfer learning in large data centers of heterogeneous disk systems. *IEEE Trans Parall Distrib Syst*, 31(9):2155-2169. <https://doi.org/10.1109/TPDS.2020.2985346>
- Zhang YQ, Hao WW, Niu B, et al., 2023. Multi-view feature-based SSD failure prediction: what, when, and why. Proc 21st USENIX Conf on File and Storage Technologies, p.409-424.
- Zhou Y, Wang F, Feng D, 2022. A disk failure prediction method based on active semi-supervised learning. *ACM Trans Stor*, 18(4):1-33. <https://doi.org/10.1145/3523699>
- Zhu BP, Wang G, Liu XG, et al., 2013. Proactive drive failure prediction for large scale storage systems. *IEEE 29th Symp on Mass Storage Systems and Technologies*, p.1-5. <https://doi.org/10.1109/MSST.2013.6558427>