

Towards Low-Energy Adaptive Personalization for Resource-Constrained Devices

Yushan Huang
Imperial College London
London, UK
yushan.huang21@imperial.ac.uk

Josh Millar
Imperial College London
London, UK
joshua.millar22@imperial.ac.uk

Yuxuan Long
Imperial College London
London, UK
maverick.long22@imperial.ac.uk

Yuchen Zhao
University of York
York, UK
yuchen.zhao@york.ac.uk

Hamed Haddadi
Imperial College London
London, UK
h.haddadi@imperial.ac.uk

Abstract

The personalization of machine learning (ML) models to address data drift is a significant challenge in the context of Internet of Things (IoT) applications. Presently, most approaches focus on fine-tuning either the full base model or its last few layers to adapt to new data, while often neglecting energy costs. However, various types of data drift exist, and fine-tuning the full base model or the last few layers may not result in optimal performance in certain scenarios. We propose Target Block Fine-Tuning (TBFT), a low-energy adaptive personalization framework designed for resource-constrained devices. We categorize data drift and personalization into three types: input-level, feature-level, and output-level. For each type, we fine-tune different blocks of the model to achieve optimal performance with reduced energy costs. Specifically, input-, feature-, and output-level correspond to fine-tuning the front, middle, and rear blocks of the model. We evaluate TBFT on a ResNet model, three datasets, three different training sizes, and a Raspberry Pi. Compared with the *BlockAvg*, where each block is fine-tuned individually and their performance improvements are averaged, TBFT exhibits an improvement in model accuracy by an average of 15.30% whilst saving 41.57% energy consumption on average compared with full fine-tuning.

CCS Concepts: • Computing methodologies → Artificial intelligence; • Computer systems organization → Embedded systems.

Keywords: Machine Learning, Resource-Constrained Devices, Low-Energy, Personalization

Permission to make digital or hard copies of part or all of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for third-party components of this work must be honored. For all other uses, contact the owner/author(s).

EuroMLSys '24, April 22, 2024, Athens, Greece

© 2024 Copyright held by the owner/author(s).

ACM ISBN XXX-X-XXXX-XXXX-X/XX/XX.

<https://doi.org/XX.XXXX/XXXXXXXX.XXXXXXX>

ACM Reference Format:

Yushan Huang, Josh Millar, Yuxuan Long, Yuchen Zhao, and Hamed Haddadi. 2024. Towards Low-Energy Adaptive Personalization for Resource-Constrained Devices. In *The 4th Workshop on Machine Learning and Systems (EuroMLSys '24)*, April 22, 2024, Athens, Greece. ACM, New York, NY, USA, 8 pages. <https://doi.org/XX.XXXX/XXXXXXXX.XXXXXXX>

1 Introduction

The deployment of deep neural networks on the Internet of Things (IoT) and mobile devices has attracted attention from both industry and academia [1–5]. In comparison to traditional cloud-based machine learning (ML), on-device ML is often susceptible to personalization issues [6, 7], where data drift arises between the source (i.e., training) and target (i.e., testing) domains, resulting in a degradation in local model performance. Additionally, due to resource constraints, ML on IoT/mobile devices imposes more stringent requirements in terms of memory, power, and energy consumption [8–10].

Several approaches have been proposed to address these issues. One approach is to enhance the domain generalization ability of the base model on the source data [11, 12]. Another approach is to fine-tune the base model on labeled target data to efficiently improve the performance [13–15]. The second approach can outperform domain generalization and unsupervised adaptation approaches in terms of model performance such as accuracy [16, 17].

Currently, most fine-tuning approaches involve fine-tuning a full model or its last few layers [18–20]. For example, in the domain of large models, a model’s front end is considered as a general feature extractor [21]. When facing new tasks, only the classifier, typically a fully connected layer, is fine-tuned. However, various types of data drift and personalization issues exist. Liang et al. [22] categorize data drift into three types: input-, feature-, and output-level. According to the Independent Causal Mechanisms (ICM) principle [23, 24], the causal generative processes of system variables consist of autonomous modules that do not inform or influence one another. From this perspective, different data shifts should correspond to local changes in the causal generative process. In other words, for various types of data

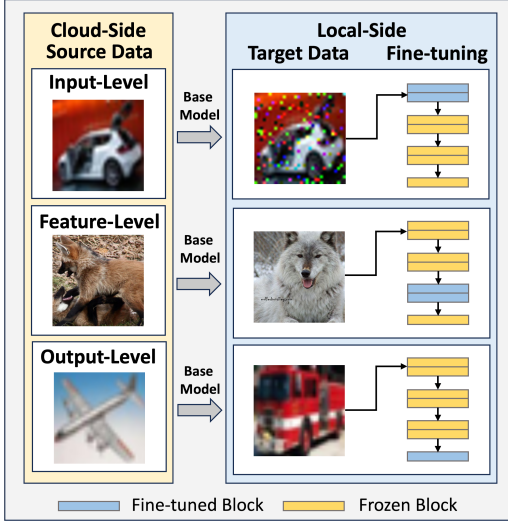


Figure 1. System Overview.

drift, fine-tuning the full model or the last few layers may not always achieve the optimal model performance. This has inspired us to explore whether fine-tuning the blocks corresponding to the different types of data drift can achieve better model performance. Furthermore, full model fine-tuning costs more energy and resources, while parameter selection methods that incorporate parameter selection into training and fine-tuning also impose additional burdens on resource-constrained devices [25–27]. Therefore, it is necessary to explore a low-energy fine-tuning and personalization method that can ensure model performance while reducing energy consumption.

Contributions.

We introduce Target Block Fine-Tuning (TBFT), a low-energy adaptive personalization framework designed for resource-constrained devices. Specifically, we categorize data drift to input-, feature-, and output-level type, TBFT fine-tunes the front block(s), middle block(s), and rear block(s) to obtain better model performance. The ‘data drift’ means the change in statistical properties of the source data compared to the target data [28]. The ‘block’ means a building unit that may encompass multiple layers such as several convolutional and pooling layers, or a single layer such as a fully connected (FC) layer. We evaluate TBFT on a ResNet model, three datasets, three training sizes, and a Raspberry Pi. Our initial experimental results show that: (i) Compared with the *BlockAvg*, where each block (Block1, Block2, Block3, and FC) is fine-tuned individually and their performance improvements are averaged, TBFT exhibits an improvement in model accuracy by an average of 15.30%; (ii) Compared with the full model fine-tuning, TBFT can save 41.57% energy consumption on average.

The source code, models, and links to the public datasets are made available on a GitHub repository ¹.

¹<https://github.com/yushan-huang/AdaptivePersonalization>

2 Background

Transfer Learning. Transfer learning (TL) facilitates the adaptation of pre-trained models to changes in their target domains. Various works have shown that freezing certain layers/parameters during fine-tuning can alleviate overfitting [29–32], as well as promote more efficient adaptation [33–35], resulting in improved TL performance in diverse settings. However, despite extensive efforts to analyze the contributions of model layers towards adaptation for the various types of data drift [36, 37], most approaches pre-select fixed layers for all types, usually the last few layers, lacking adaptability at run-time [38]. However, TBFT does not fine-tune fixed layers for all types of drift. Instead, it is more flexible and adaptive, fine-tuning different blocks for different types, thereby achieving optimal performance.

Personalization. One significant issue when deploying ML models in the wild is the personalization of pre-trained features toward shifts in target domain. Extensive research has addressed model robustness to real-world domain shifts and adaptation strategies [39–43]. It has been shown that fine-tuning the final layer is often enough for output-level shifts in domains with spurious correlations [44]. Liang et al. [22] categorized data drift into input-, feature-, and output-levels, while Lee et al. [38] found varying sensitivity to these types of drift across different model blocks, prompting us to explore block-specific fine-tuning for optimal performance across drift types. We also explore the approach’s stability on small local datasets.

Efficient Learning. There has been significant work on developing efficient ML models, reducing their size and inference energy cost, for applications on resource-constrained IoT/mobile devices with limited energy availability [45, 46]. For example, techniques such as pruning, quantization, knowledge distillation, and energy-aware architecture search [47–49] have been proposed. In this paper, we explore efficient learning from another perspective: TBFT only requires adjusting a sub-block of the model to achieve sufficient performance, which reduces the energy cost for personalization. TBFT can also be combined with the various existing efficient learning techniques, such as quantization.

3 Adaptive Personalization

Assumptions. We focus on image data and present following scenarios: (i) The base model is trained on the cloud and source datasets, the local datasets are subject to various data drifts with limited quantity. (ii) The local devices are sensitive to energy (e.g., energy-harvesting devices) [50].

We categorize data drift into three types: input-, feature-, and output-levels [22, 38]. Input-level drift refers to interference and shifts in pixel-level features in target data compared to source data, such as changes in lighting and noise interference. Feature-level implies that the source and target domains share the same categories but have different sub-populations. Output-level implies differences in output labels

between the source and target domains; for example, an image is marked as class 0 in the source data but is marked as class 1 in the target data. We believe that for different types of data drift and personalization, only fine-tuning the corresponding sub-block while freezing the other blocks can achieve sufficient performance.

In this study, we propose an adaptive block-based fine-tuning approach, TBFT, as shown in Fig. 1. On the cloud-side, we obtain a base model on the source data, then on the local-side, TBFT fine-tunes the sub-block(s) corresponding to the drift type to achieve optimal performance. We envision our model f , as a composition of function blocks f_i , represented as $f(x) = f_n \circ \dots \circ f_1(x)$, where each block f_i is equipped with a parameter set θ_i and may contain several layers. Within the TBFT, we calibrate the parameter set $\{\theta_i | i \in S\}$ from our selected block S , aiming to minimize the loss function $\hat{L}_{tgt}$. The objective is delineated as the optimization problem:

$$\min_{\theta_i, i \in S} \hat{L}_{tgt}(f(\theta_1, \dots, \theta_n)), \quad (1)$$

Here, the set S encompasses the indices of the blocks chosen for fine-tuning. Parameters θ_i not indexed in S are frozen at their pre-trained values. The block selection depends on the data drift type. For example, we fine-tune the front blocks (e.g., $S = 1$) to capture specific pixel-level variations, or the final block (e.g., $S = n$) to adapt to output-level drift. Consider fine-tuning the first block $S = 1$; it can be illustrated in the reformed model f^{tgt} , amalgamating the fine-tuned first block f_1^{tgt} with the subsequent original blocks [38], that:

$$f^{tgt}(x) = f_n \circ \dots \circ f_2 \circ f_1^{tgt}(x). \quad (2)$$

From a mathematical and logical reasoning perspective, for input-level drift, perturbations such as changes in lighting can be mitigated by fine-tuning the front block. These perturbations are akin to mapping the original input data x_{src} to a new space, producing the perturbed data $x_{trg} = Ax_{src}$ through a matrix A . This strategy is based on the intuition that the front blocks, which directly interact with the input data, are best suited to adapt to these perturbations. Fine-tuning the front blocks is essentially to directly fit and mitigate these perturbations, whereas adjusting the middle and rear blocks might not achieve optimal performance due to potential information loss in the front block processing.

Similarly, for output-level drift, where there is a change in the target data labels, fine-tuning the rear blocks to mitigate label perturbations can correct the model's predictions. Feature-level drift involves changes in the deep feature distributions of the data, where class labels remain the same, and the data distributions for the same labels in source and target domains may exhibit significant differences. Intuitively, feature-level can be viewed as the case between input- and output-levels, representing a deeper level of input-level perturbation that results in substantial feature distribution differences between the source and target domains, sufficient

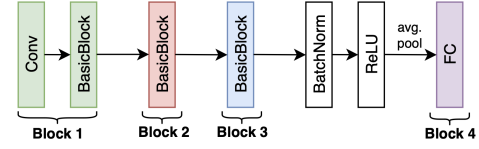


Figure 2. The ResNet-26 model used for all experiments.

to impact the model's performance. Therefore, for feature-level, fine-tuning the parameter sets of the middle blocks may compensate for the differences in feature representation between the source and target data.

4 Implementation & Evaluation Setup

In this section, we first introduce the experimental model and datasets used (Section 4.1), then detail our implementation and performance metrics (Section 4.2).

4.1 Model and Datasets

We employ the ResNet-26 [51] model, as shown in Fig. 2. We select this model due to its inherent residual block structure, enabling testing with TBFT without the need for architectural adaptation. For each experiment, we train the model from scratch on the source data. We use the Adam [52] optimizer and early-stop based on accuracy on a held-out validation set. When fine-tuning on the target data, we set a learning rate of 0.01 for the FC layer and 0.001 for the other blocks. The remaining hyperparameters vary depending on the specific experiment (details can be found in Section 5.4).

We test TBFT on three datasets, one for each category of shift (i.e. input-, feature-, and output-level). To simulate having limited local data, we set the training size to 10%, 20%, and 30% of the full dataset. The validation and testing sizes are both set at 10%.

Input-Level: Cifar10/Cifar10-C. The source distribution, drawn from the original Cifar10 dataset, consists of 32x32 color images with 10 classes. Cifar10-C, the target distribution, consists of corrupted images from the same classes at 5 different levels of severity, representing an input-level shift. There are 14 corruption types (such as frost, Gaussian blur, and Gaussian noise). For our experiments, we train and evaluate on each of the corruption types individually and report the results averaged across all types. We use the data loading setup from [53], and train/evaluate only on examples of the highest severity (i.e., corruption level 5).

Feature-Level: Living17. This dataset was created as part of the BREEDS benchmark for sub-population shift [54]. In BREEDS, ImageNet sub-populations are delineated using a WordNet hierarchy and leveraged to generate various sub-population datasets. Living17 specifically consists of images of living objects across 17 classes, each with 4 subclasses. The source and target distributions consist of different subclasses. This represents a feature-level shift, since the distribution of features varies between the source and target domains.

Output-Level: Cifar10/Cifar-Flip. To represent an output-level shift, we flip Cifar10 class labels so that each label y

Table 1. Fine-tuning accuracy results on noised blocks. The best block-based accuracy is highlighted.

		Added Noise			
		Block 1	Block 2	Block 3	FC
T u n e d	Block 1	84.88±0.23	69.24±0.40	43.44±0.38	36.30±0.41
	Block 2	72.38±0.37	83.74±0.21	70.54±0.32	67.86±0.35
	Block 3	56.86±0.37	74.95±0.27	83.04±0.29	83.85±0.39
	FC	30.58±0.24	46.94±0.23	51.82±0.41	84.28±0.33
	Block Avg	61.18±0.30	68.72±0.28	62.21±0.35	68.07±0.37
	Full	83.56±0.28	82.87±0.34	82.14±0.34	83.58±0.42

Table 2. Overall results of all experiments, including accuracy, energy cost (E), and energy-saving rate (ES) compared to full model fine-tuning. $I - L$, $F - L$, and $O - L$ correspond to input-, feature-, and output-level drift, and *BlockAvg* to the average accuracy of block-based tuning (excluding full model tuning). The best accuracy achieved on each drift is highlighted.

	Block1	Block2	Block3	FC	Block Avg	Full
I-L Acc (%)	75.20	70.98	70.62	68.74	71.40	75.09
F-L Acc (%)	48.63	57.61	68.31	60.06	58.65	67.31
O-L Acc (%)	21.25	38.69	64.27	84.74	52.29	61.11
E (J)	0.898	0.708	0.564	0.325	0.624	1.068
ES (%)	15.92	33.71	47.19	69.57	41.57	-

becomes $9 - y$ [38]. For example, label 0 is 9, and so on. The source distribution is the original Cifar10 training set.

4.2 Prototype and Performance Metrics

We evaluate TBFT on a Raspberry Pi 4 Model B with 2GB RAM, which is commonly utilized as a resource-constrained device due to its limited resource profile. The Pi 4 is equipped with a Broadcom BCM2711 processor, which is a quad-core Cortex-A72 (ARM v8) 64-bit SoC operating at 1.5GHz, offering a balance of computational power and energy efficiency. The memory is 2GB of LPDDR4-3200 SDRAM.

We evaluate the accuracy and efficiency of TBFT on various training sizes. For efficiency, We utilize the Power Monitoring HAT for Raspberry Pi to measure training system costs: (i) Runtime (s), derived from the internal clock. (ii) Memory Usage (GB), monitored internally. (iii) Energy Cost (J), determined by calculating $E = Pt$, where P is power and t is runtime. We also calculate Energy Saving (ES) by $ES = (EB - EF) / EF$, where EB and EF are the energy costs of block-based fine-tuning and full model fine-tuning, respectively. Since the experimental datasets include resolutions of 32x32 (Cifar10-C/Cifar-Flip) and 128x128 (Living17), we conduct measurements with these two resolution conditions; Since our experimental datasets have two different output sizes, 10 (Cifar10-C/Cifar-Flip) and 17 (Living17), we also conduct measurements with corresponding output sizes. Due to memory constraints, we set the batch size to 1 when evaluate the system cost on the Raspberry Pi.

5 Preliminary Evaluation Results

In this section, we evaluate TBFT on model performance and system costs.

Table 3. Training accuracy results with various batch sizes (on the Living17 dataset with 10% training data).

Batch Size	Block 1	Block 2	Block 3	FC	Full
128	47.30	55.57	65.83	59.01	64.69
1	45.17	54.12	64.23	57.92	62.97

5.1 Motivation Experiments

To validate whether adjusting different model blocks is sufficient for model personalization towards different data drifts, we introduce normally distributed noise to various model blocks, including the last FC layer. This simulates distribution shifts specifically localized to those parameters. We do this on a ResNet-26 model [51], trained from scratch on Cifar-10 [55], fine-tuning the different blocks of the network on Cifar-10 target data while freezing all other parameters.

Table 1 shows the results of our motivation experiments; these results show that selectively fine-tuning just the noisy block outperforms fine-tuning other blocks and rivals or exceeds full model fine-tuning, with an average accuracy increase of 18.94% over *BlockAvg*. This indicates targeted block fine-tuning can achieve optimal personalization.

5.2 Overall Performance

Table. 2 provides a summary of all experimental results, including accuracy (across three types of drift and with multiple training sizes) and system cost (across all training configurations). We find that fine-tuning Block 1, Block 3, and the FC layer, corresponding to input-, feature-, and output-level drifts, can achieve performance approaching or even surpassing full model fine-tuning. Specifically, compared with *BlockAvg*, TBFT can improve model accuracy by 15.30% on average, suggesting a link between different drift types and the most affected model blocks, where fine-tuning the relevant block can markedly improve performance. In addition, compared with full model fine-tuning, TBFT can save 41.53% energy cost on average. Table. 2 summarizes the results from all experiments, with further details in the following section.

5.3 Single-Batch & Multi-Batch Training

When training on resource-constrained devices, the batch size is generally set to one [56]. However, this can also lead to a decrease in the efficiency of experiments. Therefore, at the beginning of our experiments, we conduct both single-batch training (batch size = 1) on the Raspberry Pi and multi-batch training (batch size = 128) on the GPU, and compare their performance to verify if the latter can be used as an approximate method to increase experiment efficiency. All other training settings are kept consistent. Due to the lower efficiency of single-batch training, we conduct tests only on the Living17 [54] dataset, selecting only 10% of the data as training data. The results are shown in the Table. 3.

We find that single-batch training and multi-batch training have similar accuracy under the same training configuration. This allows us to use the results from multi-batch training on the GPU for evaluation to increase experiment efficiency.

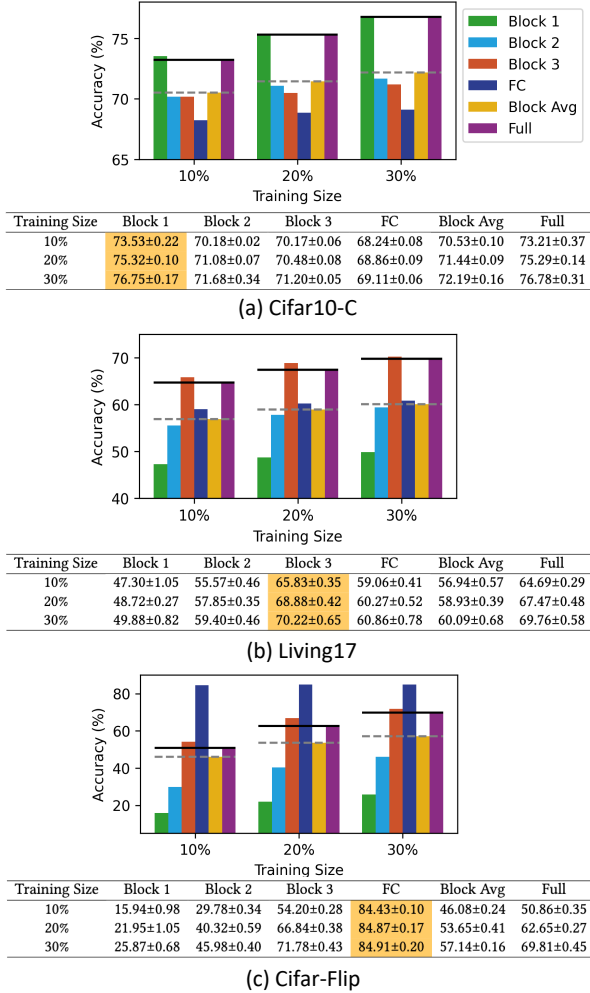


Figure 3. The accuracy results of TBFT on three datasets and with different training sizes. The horizontal solid line shows the accuracy of *BlockAvg*, and the horizontal dashed line the accuracy of full model fine-tuning. The highlighted values are the best accuracy across block-based fine-tuning.

Therefore, unless stated otherwise, we default to reporting model performance based on GPU training.

5.4 Model Performance

We test the model personalization accuracy for different types of data drift on three datasets: Cifar10-C [57] (input-level), Living17 [54] (feature-level), and Cifar-Flip [58] (output-level), using the ResNet-26 [51]. To evaluate the accuracy of the proposed method on small local datasets, we test with 10%, 20%, and 30% training sizes.

Input-Level: Cifar10-C. We load the ResNet-26 [51] model and train from scratch on the Cifar10 dataset (with 10 output classes and image resolution of 32x32). Subsequently, we conduct block fine-tuning on the Cifar10-C dataset, activating only one block and freezing the others. The results are shown in Fig. 3 (a). We find that when processing input-level drift, fine-tuning only Block 1 achieves the best overall

model performance, approaching or even surpassing full model fine-tuning. Specifically, compared to *BlockAvg*, fine-tuning only Block 1 at training sizes of 10%, 20%, and 30% results in average improvements in model accuracy of 3.00%, 3.88%, and 4.56%, respectively. Moreover, the closer a block is to the front of the model, the better its performance is. The best model performance is achieved by fine-tuning the front-most Block 1 while the worst performance is from fine-tuning the FC layer. Fine-tuning the middle blocks (Blocks 2 and 3) results in moderate performance. These experimental results align with our analysis in Section. 3, indicating that for input-level drift, the front block of the model is more sensitive than the other blocks.

Feature-Level: Living17. We train our ResNet-26 base model from scratch on the source dataset from Living17 (with 17 output classes and image resolution of 128x128). Our data preprocessing and augmentation procedures remain consistent with those used in BREEDS [54], which introduced the Living17 dataset. The results are shown in Fig. 3 (b). We find that for feature-level drift, fine-tuning Block 3 results in the best overall performance, slightly outperforming full model fine-tuning. Compared to *BlockAvg*, fine-tuning Block 3 at training sizes of 10%, 20%, and 30% results in average accuracy improvements of 8.89%, 9.95%, and 10.13%, respectively. Moreover, the results indicate that the closer the block to the middle of model, the better the model performance is. For instance, at a training size of 30%, with Block 3 serving as a reference point (achieving an accuracy of 70.22%), fine-tuning the adjacent blocks (i.e., Block 2 and the FC layer) results in model accuracy of 59.40% and 60.86%, respectively. However, fine-tuning Block 1, which is the farthest from Block 3, yields an accuracy of only 49.88%. These findings underscore the critical role of the intermediate blocks in capturing and adapting to feature-level drift.

Output-Level: Cifar-Flip. The model setup, training, and fine-tuning procedures are the same as in the Cifar10-C experiments, with the only difference being Cifar10-C is replaced with Cifar-Flip. The results are shown in Fig. 3 (c). We find that, similar to the previous experimental results, for output-level drift, fine-tuning the blocks closer to the model's rear results in better accuracy, with the best accuracy being achieved through fine-tuning the FC layer. Specifically, compared to *BlockAvg*, fine-tuning the FC layer at training sizes of 10%, 20%, and 30% results in accuracy improvements of 38.35%, 31.22%, and 27.77%, respectively. This surpasses fine-tuning the full model, which shows accuracy improvements of 33.57%, 22.22%, and 15.10%, respectively. We also find that fine-tuning blocks closer to the model's rear results in higher model accuracy. For instance, at a training size of 30%, fine-tuning the block from the front to back results in accuracies of 25.87%, 45.98%, 71.78%, and 84.91%. Additionally, we can also observe that fine-tuning only the FC layer requires only 10% training size to achieve performance close to that with 30% training size. This indicates that for output-level data

Table 4. The time and energy costs of block-based and full model fine-tuning. The *Energy – SavingRate* is calculated by comparing the current energy cost to the energy cost of full model fine-tuning.

	Resolution 32, Output 10						Resolution 128, Output 17					
	Block 1	Block 2	Block 3	FC	Block Avg	All	Block 1	Block 2	Block 3	FC	Block Avg	All
Time (s)	0.45	0.33	0.29	0.17	0.31	0.53	6.82	5.40	4.28	2.46	4.74	8.12
Energy Cost (J)	0.090	0.066	0.058	0.034	0.062	0.106	1.705	1.350	1.070	0.615	1.185	2.030
Energy-Saving Rate (%)	15.09	37.74	45.28	67.92	41.51	-	16.01	33.50	47.29	69.70	41.63	-

drift, it is possible to fine-tune the FC layer with a smaller local data size and achieve sufficient model performance.

Summary. Our results indicate that the best model block to fine-tune varies for different types of data drift. Specifically, for input-, feature-, and output-level, the best blocks to fine-tune are the front block (e.g., Block 1), the middle block(s) (e.g., Block 3), and the rear block (e.g., the FC layer). For the scenario where one aspect of the distribution changes while others remain unchanged, it is only necessary to fine-tune the corresponding block to achieve sufficient performance. For input-level drift (e.g., image interference), the pixel-level features of the image change, but the basic structure of the data remains the same between the source and target data. Therefore, it is only necessary to fine-tune the front block to adapt to this image interference. For output-level drift (e.g., label shift), the pixel-level features of the image do not change, and only the output labels have drifted, thus it is only necessary to fine-tune the rear block. Feature-level drift can be considered as somewhere between input- and output-level, where not only is there a significant change in image distribution (i.e. image interference), but this change also brings about an output-level disturbance. Therefore, it is best to fine-tune the middle block(s) of the model. These preliminary results demonstrate that TBFT is a viable way for low-energy personalization in resource-constrained devices.

5.5 System Cost

We measure the system cost on a Raspberry Pi by Power Monitoring HAT. The input voltage is 5V and standby memory usage is 0.75GB. In both TBFT and full training, with resolution 32 and output size 10, the operational memory usage remains at approximately 0.11GB. Similarly, under resolution 128 and output size 17, the operational memory usage is about 0.28GB. The memory usage meets the requirements of the experimental device.

The standby power consumption is about 4W. The operational power consumption is about 4.20W with resolution 32x32 and output size 10, and about 4.25W with resolution 128x128 and output size 17. The energy cost and time under different configurations are shown in Table. 4. We find that TBFT results in varying degrees of energy savings. Compared to full model fine-tuning, TBFT can achieve a minimum energy saving of 15.09%, a maximum of 69.70%, and an average of 41.57%. Furthermore, across different training configurations, TBFT consistently demonstrates stable energy savings. Overall, analyzing the experimental results in conjunction

with model performance and system cost, we find that TBFT can achieve better model performance at lower energy costs.

6 Discussion and Future Work

Key Findings. Our preliminary results indicate that fine-tuning the blocks corresponding to drift types, i.e., TBFT, can achieve better model performance and significantly reduce power consumption. This lays the foundation for achieving low-power personalization on resource-constrained devices. Future work will focus on the following aspects:

Discrimination Head. Currently, TBFT relies on prior knowledge of the type of data drift. However, in practical applications, the type of drift is often difficult to know in advance. Some parameter/block selection methods are based on the gradient changes during model training, such as Relative Gradient Norm (Auto-RGN) and Signal-to-Noise Ratio (Auto-SNR) [38]. However, these methods embed the parameter/module selection process into the entire training process, which is an additional burden for resource-constrained devices and may affect performance. We plan to approach parameter/block selection from the perspective of the data itself, without embedding it into the training process. For example, we could use information theory [59, 60] or contrastive learning [61, 62] methods to select the parameters/block that need to be trained before model training based on the characteristics of the dataset.

Unsupervised Personalization. We also plan to explore unsupervised personalization methods. Zhang et al. propose MEMO, offering a novel method for unsupervised learning by minimizing the marginal entropy of average predictions for individual images [51]. Considering the implementation principle of MEMO, it may yield good results on input- and feature-level drift. However, for output-level drift, MEMO's strategy of minimizing marginal entropy may backfire, as this optimization direction may conflict with the true distribution of new labels, thereby reducing model performance. In this case, a clustering-based classifier provides a potential solution [63]. The core advantage of this approach lies in its reliance on data feature representations rather than label information, enabling it to remain effective even in cases of label inversion. Additionally, the clustering-based classifier may also be effective for input- and feature-level data drifts.

Multidimensional Personalization. In real applications, the data often faces not isolated but complex and composite types of drift, which may simultaneously be affected by input-, feature-, and output-level drifts. This multidimensional data

drift necessitates the development of a personalization framework capable of comprehensively addressing these complex scenarios. One possible solution is to first detect and rank the importance of different types of data drift using the Discrimination Head. Subsequently, considering the resource constraints of IoT devices, we adopt a segmented and joint optimization strategy to gradually adapt to different types of data drift, for example, by leveraging the principles of early exit mechanisms [64].

7 Conclusion

In this paper, we introduced TBFT, a practical low-energy framework to address the challenges of data drift and model personalization on resource-constrained devices. TBFT categorizes data drift into three types: input-, feature-, and output-level. For the different types, fine-tuning their corresponding blocks of the model can achieve the best model performance with reduced energy cost. Our preliminary experimental results show that TBFT exhibits an improvement in accuracy by an average of 15.30%, and can save 41.57% energy consumption on average. TBFT's adaptability, energy efficiency, and performance make it a feasible solution for low-energy adaptive personalization on mobile devices.

References

- [1] MG Sarwar Murshed, Christopher Murphy, Daqing Hou, Nazar Khan, Ganesh Ananthanarayanan, and Faraz Hussain. Machine learning at the network edge: A survey. *ACM Computing Surveys (CSUR)*, 54(8):1–37, 2021.
- [2] Ke Zhang, Hengchang Liu, and Siobhán Clarke. Dbgan: A data balancing generative adversarial network for mobility pattern recognition. In *International Conference on Big Data Analytics and Knowledge Discovery*, pages 120–134. Springer, 2023.
- [3] Yushan Huang, Yu Chen, and Zhenyu Hu. Thermal error modeling and analysis of CNC machine tools based on wavelet neural network. In *2021 IEEE International Conference on Consumer Electronics and Computer Engineering (ICCECE)*, pages 454–457. IEEE, 2021.
- [4] Liangyi Gong, Zhenhua Li, Feng Qian, Zifan Zhang, Qi Alfred Chen, Zhiyun Qian, Hao Lin, and Yunhao Liu. Experiences of landing machine learning onto market-scale mobile malware detection. In *Proceedings of the Fifteenth European Conference on Computer Systems*, pages 1–14, 2020.
- [5] Ji-Ying Shi, Le-Tao Ling, Fei Xue, Zi-Jian Qin, Ya-Jing Li, Zhi-Xin Lai, and Ting Yang. Combining incremental conductance and firefly algorithm for tracking the global mpp of pv arrays. *Journal of Renewable and Sustainable Energy*, 9(2), 2017.
- [6] Ozlem Durmaz Incel and Sevda Ooze Bursa. On-device deep learning for mobile and wearable sensing applications: A review. *IEEE Sensors Journal*, 2023.
- [7] Saupatik Dhar, Junyao Guo, Jiayi Liu, Samarth Tripathi, Unmesh Kurup, and Mohak Shah. A survey of on-device machine learning: An algorithms and learning theory perspective. *ACM Transactions on Internet of Things*, 2(3):1–49, 2021.
- [8] Yushan Huang and Hamed Haddadi. Poster: Towards Battery-Free Machine Learning Inference and Model Personalization on MCUs. In *Proceedings of the 21st Annual International Conference on Mobile Systems, Applications and Services*, pages 571–572, 2023.
- [9] Yuchen Zhao, Sayed Saad Afzal, Waleed Akbar, Osvy Rodriguez, Fan Mo, David Boyle, Fadel Adib, and Hamed Haddadi. Towards battery-free machine learning and inference in underwater environments. In *Proceedings of the 23rd Annual International Workshop on Mobile Computing Systems and Applications*, pages 29–34, 2022.
- [10] Sayed Saad Afzal, Waleed Akbar, Osvy Rodriguez, Mario Doumet, Unsoo Ha, Reza Ghaffarivardavagh, and Fadel Adib. Battery-free wireless imaging of underwater environments. *Nature communications*, 13(1):5546, 2022.
- [11] Bernhard Schölkopf, Francesco Locatello, Stefan Bauer, Nan Rosemary Ke, Nal Kalchbrenner, Anirudh Goyal, and Yoshua Bengio. Toward causal representation learning. *Proceedings of the IEEE*, 109(5):612–634, 2021.
- [12] Junbum Cha, Sanghyuk Chun, Kyungjae Lee, Han-Cheol Cho, Seunghyun Park, Yunsung Lee, and Sungrae Park. Swad: Domain generalization by seeking flat minima. *Advances in Neural Information Processing Systems*, 34:22405–22418, 2021.
- [13] Nattaya Mairittha, Tittaya Mairittha, and Sozo Inoue. Improving activity data collection with on-device personalization using fine-tuning. In *Adjunct Proceedings of the 2020 ACM International Joint Conference on Pervasive and Ubiquitous Computing and Proceedings of the 2020 ACM International Symposium on Wearable Computers*, pages 255–260, 2020.
- [14] Hong-You Chen, Yandong Li, Yin Cui, Mingda Zhang, Wei-Lun Chao, and Li Zhang. Train-Once-for-All Personalization. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 11818–11827, 2023.
- [15] Zhixin Lai, Xuesheng Zhang, and Suiyao Chen. Adaptive ensembles of fine-tuned transformers for llm-generated text detection. *arXiv preprint arXiv:2403.13335*, 2024.
- [16] Elan Rosenfeld, Pradeep Ravikumar, and Andrej Risteski. Domain-adjusted regression or: Erm may already learn features sufficient for out-of-distribution generalization. *arXiv preprint arXiv:2202.06856*, 2022.
- [17] Polina Kirichenko, Pavel Izmailov, and Andrew Gordon Wilson. Last Layer Re-Training is Sufficient for Robustness to Spurious Correlations. In *ICML 2022: Workshop on Spurious Correlations, Invariance and Stability*, 2022.
- [18] Mitchell Wortsman, Gabriel Ilharco, Jong Wook Kim, Mike Li, Simon Kornblith, Rebecca Roelofs, Raphael Gontijo Lopes, Hannaneh Hajishirzi, Ali Farhadi, Hongseok Namkoong, et al. Robust fine-tuning of zero-shot models. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 7959–7971, 2022.
- [19] Zhi Wang, Wei Bi, Yan Wang, and Xiaojiang Liu. Better fine-tuning via instance weighting for text classification. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 33, pages 7241–7248, 2019.
- [20] Yushan Huang, Ranya Aloufi, Xavier Cadet, Yuchen Zhao, Payam Barnaghi, and Hamed Haddadi. Microt: Low-energy and adaptive models for mcus. *arXiv preprint arXiv:2403.08040*, 2024.
- [21] Maxime Oquab, Timothée Darcet, Théo Moutakanni, Huy Vo, Marc Szafraniec, Vasil Khalidov, Pierre Fernandez, Daniel Haziza, Francisco Massa, Alaaeldin El-Nouby, et al. Dinov2: Learning robust visual features without supervision. *arXiv preprint arXiv:2304.07193*, 2023.
- [22] Jian Liang, Ran He, and Tieniu Tan. A comprehensive survey on test-time adaptation under distribution shifts. *arXiv preprint arXiv:2303.15361*, 2023.
- [23] Bernhard Schölkopf, Dominik Janzing, Jonas Peters, Eleni Sgouritsa, Kun Zhang, and Joris Mooij. On causal and anticausal learning. *arXiv preprint arXiv:1206.6471*, 2012.
- [24] Jonas Peters, Dominik Janzing, and Bernhard Schölkopf. *Elements of causal inference: foundations and learning algorithms*. The MIT Press, 2017.
- [25] Elad Ben Zaken, Shauli Ravfogel, and Yoav Goldberg. Bitfit: Simple parameter-efficient fine-tuning for transformer-based masked language-models. *arXiv preprint arXiv:2106.10199*, 2021.
- [26] Ning Ding, Yujia Qin, Guang Yang, Fuchao Wei, Zonghan Yang, Yusheng Su, Shengding Hu, Yulin Chen, Chi-Min Chan, Weize Chen,

- et al. Parameter-efficient fine-tuning of large-scale pre-trained language models. *Nature Machine Intelligence*, 5(3):220–235, 2023.
- [27] Young D Kwon, Rui Li, Stylianos I Venieris, Jagmohan Chauhan, Nicholas D Lane, and Cecilia Mascolo. Tinytrain: Deep neural network training at the extreme edge. *arXiv preprint arXiv:2307.09988*, 2023.
- [28] Adriana Sayuri Iwashita and Joao Paulo Papa. An overview on concept drift learning. *IEEE access*, 7:1532–1547, 2018.
- [29] James Kirkpatrick, Razvan Pascanu, Neil C. Rabinowitz, Joel Veness, Guillaume Desjardins, Andrei A. Rusu, Kieran Milan, John Quan, Tiago Ramalho, Agnieszka Grabska-Barwinska, Demis Hassabis, Claudia Clopath, Dharshan Kumaran, and Raia Hadsell. Overcoming catastrophic forgetting in neural networks. *CoRR*, abs/1612.00796, 2016.
- [30] Jaejun Lee, Raphael Tang, and Jimmy Lin. What Would Elsa Do? Freezing Layers During Transformer Fine-Tuning. *CoRR*, abs/1911.03090, 2019.
- [31] Yunhui Guo, Honghui Shi, Abhishek Kumar, Kristen Grauman, Tajana Rosing, and Rogério Schmidt Feris. SpotTune: Transfer Learning through Adaptive Fine-tuning. *CoRR*, abs/1811.08737, 2018.
- [32] Vinay V. Ramasesh, Ethan Dyer, and Maithra Raghu. Anatomy of Catastrophic Forgetting: Hidden Representations and Task Semantics. *CoRR*, abs/2007.07400, 2020.
- [33] Yuhao Liu, Saurabh Agarwal, and Shivaram Venkataraman. AutoFreeze: Automatically Freezing Model Blocks to Accelerate Fine-tuning. *CoRR*, abs/2102.01386, 2021.
- [34] Amelie Royer and Christoph H. Lampert. A Flexible Selection Scheme for Minimum-Effort Transfer Learning. *CoRR*, abs/2008.11995, 2020.
- [35] Cian Eastwood, Ian Mason, and Christopher K. I. Williams. Unit-level surprise in neural networks. In Melanie F. Pradier, Aaron Schein, Stephanie Hyland, Francisco J. R. Ruiz, and Jessica Z. Forde, editors, *Proceedings on "I (Still) Can't Believe It's Not Better!" at NeurIPS 2021 Workshops*, volume 163 of *Proceedings of Machine Learning Research*, pages 33–40. PMLR, 13 Dec 2022.
- [36] Behnam Neyshabur, Hanie Sedghi, and Chiyuan Zhang. What is being transferred in transfer learning? *CoRR*, abs/2008.11687, 2020.
- [37] Chiyuan Zhang, Samy Bengio, and Yoram Singer. Are All Layers Created Equal?, 2022.
- [38] Yoonho Lee, Annie S. Chen, Fahim Tajwar, Ananya Kumar, Huaxiu Yao, Percy Liang, and Chelsea Finn. Surgical Fine-Tuning Improves Adaptation to Distribution Shifts, 2023.
- [39] Ananya Kumar, Aditi Raghunathan, Robbie Jones, Tengyu Ma, and Percy Liang. Fine-Tuning can Distort Pretrained Features and Underperform Out-of-Distribution, 2022.
- [40] Anders Andreassen, Yasaman Bahri, Behnam Neyshabur, and Rebecca Roelofs. The Evolution of Out-of-Distribution Robustness Throughout Fine-Tuning. *CoRR*, abs/2106.15831, 2021.
- [41] Olivia Wiles, Sven Goyal, Florian Stimberg, Sylvestre-Alvise Rebuffi, Ira Ktena, Krishnamurthy Dvijotham, and A. Taylan Cemgil. A Fine-Grained Analysis on Distribution Shift. *CoRR*, abs/2110.11328, 2021.
- [42] Sang Michael Xie, Ananya Kumar, Robbie Jones, Fereshte Khani, Tengyu Ma, and Percy Liang. In-N-Out: Pre-Training and Self-Training using Auxiliary Information for Out-of-Distribution Robustness. *CoRR*, abs/2012.04550, 2020.
- [43] Hadi Salman, Andrew Ilyas, Logan Engstrom, Ashish Kapoor, and Aleksander Madry. Do Adversarially Robust ImageNet Models Transfer Better? *CoRR*, abs/2007.08489, 2020.
- [44] Polina Kirichenko, Pavel Izmailov, and Andrew Gordon Wilson. Last Layer Re-Training is Sufficient for Robustness to Spurious Correlations, 2023.
- [45] Dr. Lachit Dutta and Swapna Bharali. TinyML Meets IoT: A Comprehensive Survey. *Internet of Things*, 16:100461, 2021.
- [46] Ramon Sanchez-Iborra and Antonio F. Skarmeta. TinyML-Enabled Frugal Smart Objects: Challenges and Opportunities. *IEEE Circuits and Systems Magazine*, 20(3):4–18, 2020.
- [47] Tien-Ju Yang, Yu-Hsin Chen, and Vivienne Sze. Designing Energy-Efficient Convolutional Neural Networks Using Energy-Aware Pruning. In *2017 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 6071–6079, 2017.
- [48] Chengyue Gong, Zixuan Jiang, Dilin Wang, Yibo Lin, Qiang Liu, and David Z. Pan. Mixed Precision Neural Architecture Search for Energy Efficient Deep Learning. In *2019 IEEE/ACM International Conference on Computer-Aided Design (ICCAD)*, pages 1–7, 2019.
- [49] Daniel T. Speckhard, Karolis Misiunas, Sagi Perel, Tenghui Zhu, Simon Carlile, and Malcolm Slaney. Neural Architecture Search for Energy Efficient Always-on Audio Models, 2023.
- [50] Seunghyeok Jeon, Yonghun Choi, Yeonwoo Cho, and Hojung Cha. HarvNet: Resource-Optimized Operation of Multi-Exit Deep Neural Networks on Energy Harvesting Devices. In *Proceedings of the 21st Annual International Conference on Mobile Systems, Applications and Services*, pages 42–55, 2023.
- [51] Marvin Zhang, Sergey Levine, and Chelsea Finn. Memo: Test time robustness via adaptation and augmentation. *Advances in Neural Information Processing Systems*, 35:38629–38642, 2022.
- [52] Diederik P. Kingma and Jimmy Ba. Adam: A method for stochastic optimization, Jan 2017.
- [53] Francesco Croce, Maksym Andriushchenko, Vikash Sehwal, Nicolas Flammarion, Mung Chiang, Prateek Mittal, and Matthias Hein. RobustBench: a standardized adversarial robustness benchmark. *CoRR*, abs/2010.09670, 2020.
- [54] Shibani Santurkar, Dimitris Tsipras, and Aleksander Madry. BREEDS: Benchmarks for Subpopulation Shift. In *ArXiv preprint arXiv:2008.04859*, 2020.
- [55] Alex Krizhevsky, Geoffrey Hinton, et al. Learning multiple layers of features from tiny images. 2009.
- [56] Ji Lin, Ligeng Zhu, Wei-Ming Chen, Wei-Chen Wang, Chuang Gan, and Song Han. On-device training under 256kb memory. *Advances in Neural Information Processing Systems*, 35:22941–22954, 2022.
- [57] Dan Hendrycks and Thomas Dietterich. Benchmarking Neural Network Robustness to Common Corruptions and Perturbations. *Proceedings of the International Conference on Learning Representations*, 2019.
- [58] Shiori Sagawa, Pang Wei Koh, Tatsunori B Hashimoto, and Percy Liang. Distributionally Robust Neural Networks. In *International Conference on Learning Representations*, 2019.
- [59] Lei Du, Qinbao Song, and Xiaolin Jia. Detecting concept drift: an information entropy based method using an adaptive sliding window. *Intelligent Data Analysis*, 18(3):337–364, 2014.
- [60] Zilu Liang, Mario Alberto Chapa Martell, and Takuichi Nishimura. A personalized approach for detecting unusual sleep from time series sleep-tracking data. In *2016 IEEE International Conference on Healthcare Informatics (ICHI)*, pages 18–23. IEEE, 2016.
- [61] Mingrui Lao, Nan Pu, Zhun Zhong, Nicu Sebe, and Michael S Lew. FedVQA: Personalized Federated Visual Question Answering over Heterogeneous Scenes. In *Proceedings of the 31st ACM International Conference on Multimedia*, pages 7796–7807, 2023.
- [62] Yupei Zhang, Yunan Xu, Shuangshuang Wei, Yifei Wang, Yuxin Li, and Xuequn Shang. Doubly contrastive representation learning for federated image recognition. *Pattern Recognition*, 139:109507, 2023.
- [63] Hao Wu, Jinghao Feng, Xuejin Tian, Edward Sun, Yunxin Liu, Bo Dong, Fengyuan Xu, and Sheng Zhong. Emo: Real-time emotion recognition from single-eye images for resource-constrained eyewear devices. In *Proceedings of the 18th International Conference on Mobile Systems, Applications, and Services*, pages 448–461, 2020.
- [64] Mary Phuong and Christoph H Lampert. Distillation-based training for multi-exit architectures. In *Proceedings of the IEEE/CVF international conference on computer vision*, pages 1355–1364, 2019.