
Evaluating Continual Test-Time Adaptation for Contextual and Semantic Domain Shifts

Tommie Kerssies
Eindhoven University of Technology
t.kerssies@student.tue.nl

Mert Kılıçkaya
Eindhoven University of Technology
m.kilickaya@tue.nl

Joaquin Vanschoren
Eindhoven University of Technology
j.vanschoren@tue.nl

Abstract

In this paper, our goal is to adapt a pre-trained convolutional neural network to domain shifts at test time. We do so continually with the incoming stream of test batches, without labels. The existing literature mostly operates on artificial shifts obtained via adversarial perturbations of a test image. Motivated by this, we evaluate the state of the art on two realistic and challenging sources of domain shifts, namely contextual and semantic shifts. Contextual shifts correspond to the environment types, for example, a model pre-trained on indoor context has to adapt to the outdoor context on CORE-50 [7]. Semantic shifts correspond to the capture types, for example a model pre-trained on natural images has to adapt to cliparts, sketches, and paintings on DomainNet [10]. We include in our analysis recent techniques such as Prediction-Time Batch Normalization (BN) [8], Test Entropy Minimization (TENT) [16] and Continual Test-Time Adaptation (CoTTA) [17]. Our findings are three-fold: *i*) Test-time adaptation methods perform better and forget less on contextual shifts compared to semantic shifts, *ii*) TENT outperforms other methods on short-term adaptation, whereas CoTTA outperforms other methods on long-term adaptation, *iii*) BN is most reliable and robust. Our code is available at <https://github.com/tommiekerssies/Evaluating-Continual-Test-Time-Adaptation-for-Contextual-and-Semantic-Domain-Shifts>.

1 Introduction

Motivation for domain adaptation. Distributional shifts in evaluation data alter a model’s performance. Domain adaptation is proposed to adapt an existing model to novel situations. More specifically, domain adaptation adapts a model trained on one or more source domains to a target domain. Regular domain adaptation considers a stationary target, meaning that the target does not change over time.

Motivation for continual learning. Continual learning is a concept to learn a model on data with continually shifting distributions or even changing tasks, without forgetting what was learned previously. Storing data may not be feasible due to storage or privacy concerns, therefore we only assume access to new data in the sequence. In a domain adaptation context, having access to the entire target dataset a priori may not be realistic, because real-world data is collected continually and the distribution may shift over time.

Continual domain adaptation. To address the need for a model that can adapt over time, continual domain adaptation is proposed. More specifically, this work focuses on continual test-time adaptation, where we do not assume access to the source data and adapt online to continually arriving unlabeled test batches.

Contextual and semantic domain shifts. This work focuses on visual recognition tasks (image classification). We distinguish two types of distribution shifts: contextual and semantic shifts. In a semantic distribution shift, the semantics of the relevant parts of the image change (e.g. an object is captured in real life instead of a drawing). In a contextual shift, the semantics of the relevant parts of the image remain the same, but may appear differently due to a context change (e.g. the lighting conditions change).

Evaluation of continual domain adaptation in contextual and semantic domain shifts. We evaluate Prediction-Time Batch Normalization (BN) [8], Test Entropy Minimization (TENT) [16] and Continual Test-Time Adaptation (CoTTA) [17] for continual test-time adaptation to contextual distribution shifts on CORE50 [7] and semantic distribution shifts on DomainNet [10]. We find that the test-time adaptation methods show relatively more improvements and less forgetting on contextual shifts compared to semantic shifts. For short-term adaptation, we find that TENT performs best and for long-term adaptation, we find that CoTTA performs best. Finally, we find that unlike TENT or CoTTA, BN does not show error accumulation or catastrophic forgetting, and works best when the source model was trained on a dataset with a relatively narrow distribution.

2 Related Work

2.1 Transfer Learning

Transfer learning is about transferring knowledge from a source to a target. Transfer learning is a very broad term and can be divided into different subsettings. Following the work of [11], we consider different scenarios. First of all, whether the source and target data are from the same distribution or not. If source and target are from the same distribution, and the task of source and target is the same, this is the usual learning setting. If source and target are from the same distribution, but their task differs, we call this inductive transfer learning. Similarly, the distributions may differ but the task may be the same, which is called transductive transfer learning. Finally, both the distribution and task may be different, called unsupervised transfer learning.

2.2 Domain Adaptation

Transductive transfer learning, where source and target distributions differ but their task is the same, is more commonly referred to as domain adaptation. Table 1 outlines related settings of domain adaptation with their similarities and differences.

Table 1: Settings of domain adaptation.

Setting	Source access	Target labels	Target stationary
Supervised domain adaptation	✓	✓	✓
Unsupervised domain adaptation	✓		✓
Test-time adaptation			✓
Continual test-time adaptation			

2.2.1 Test-Time Adaptation

Test-time adaptation is similar to unsupervised domain adaptation, where no target labels are available. However, unsupervised domain adaptation requires access to the source data when adapting to the target data. Test-time adaptation methods do not require this access, and can therefore adapt to target data at test-time. Table 2 outlines methods for test-time adaptation that we will evaluate in this work in the setting where the target is continually changing and not stationary.

Test-time adaptation methods can be divided into calibration or optimization based methods. Calibration based methods only update the batch normalization statistics [4], whereas optimization based methods will update the learnable parameters in the model with some unsupervised objective.

Calibration based. The concurrent works of [14] and [8] investigate the effect of updating the statistics at test-time to reduce covariate shift, where [8] simply proposes to discard the source statistics at test-time and only use test batch statistics. A notable difference with [14] is that [14] proposes to calibrate the batch normalization statistics by combining the source statistics with the test batch statistics using a certain formula with a hyperparameter that can emphasize the source or the target more. Other works, such as α -BN [18] propose a similar thing. Most of these methods focus mostly on using test batches, however the work of SITA [5] looks into approximating test batch statistics one test image at a time. In the work of [19], the authors propose to learn how to calibrate the batch normalization statistics using a meta-model, which should also work on one test image at a time.

Optimization based. As test batches do not have labels, an optimization based approach to test-time adaptation needs another objective to optimize for.

The authors of TENT [16] propose to update the batch normalization transformation parameters by minimizing test entropy. CoTTA [17] builds on TENT for a continual learning setting, by updating all parameters and introducing weight and augmentation averaged pseudo-labels as well as stochastic restoration. Many other works are also based on TENT and another noteworthy one is [9], where the authors propose to only minimize entropy when entropy is below a certain threshold, which should minimize error accumulation by minimizing entropy on uncertain samples.

Another approach to optimization based test-time adaptation is self-supervised learning. Test-time training as proposed in [15] utilizes multi-task learning for a supervised and a self-supervised task when training the source model, and will keep on training the self-supervised task at test-time. A distinction with the setting studied in our work is that this approach will not work on any pre-trained source model and requires the source model to be trained in a certain way. AdaContrast [1] uses a self-supervised objective based on contrastive learning for test-time adaptation, but should also work in our setting as it does not require a specifically trained source model.

Table 2: Methods for test-time adaptation.

Method	Strategy	Updated parameters	Supervision
BN [8]	Calibration	$E[x], Var[x]$	Unsupervised
TENT [16]	Calibration, Optimization	$E[x], Var[x], \gamma, \beta$	Pseudo-supervised
CoTTA [17]	Calibration, Optimization	All	Pseudo-supervised

2.3 Domain Generalization

In domain generalization terminology, the data a model was trained on is referred to as in-distribution. Data with a distribution different from the training data would be called out-of-distribution. Domain generalization methods aim to train a model to be more generalizable, that is, better performing at out-of-distribution (test) data. Different from domain adaptation, where access to the target domain in some way is required, domain generalization techniques can be applied without having access to any out-of-distribution data. Domain generalization methods can be combined with domain adaptation methods if there is also access to out-of-distribution data. In our work, we do not utilize any domain generalization methods.

3 Empirical Method

Prediction-time batch normalization (BN). Batch normalization layers in a neural network are shown to make training faster and more stable by normalizing neuron inputs [4]. The authors of [4] hypothesize that Batch Normalization layers mitigate the problem of internal covariate shift, which causes performance degradation as well as worsened uncertainty estimation.

Mathematically, a batch normalization layer can be expressed as:

$$y = \frac{x - E[x]}{\sqrt{\text{Var}[x] + \epsilon}} * \gamma + \beta \quad (1)$$

where γ and β are learnable transformation parameters and $E[x]$ and $\text{Var}[x]$ are an estimation of the mean input to the neuron and its variance. These neuron input statistics are commonly estimated during training by using an exponential moving average that is updated for every batch encountered during training:

$$\hat{x}_{t+1} = \alpha * \hat{x}_t + (1 - \alpha) * x_{t+1} \quad (2)$$

where $\hat{x}_t$ is the exponential moving average of the statistic on all previous batches and x_{t+1} is the value of the statistic on the newly encountered batch (and α is a smoothing factor). After training, the estimated neuron input statistics are frozen. Consequently, the frozen neuron input statistics are used for inference.

In case of a distribution shift at test-time, neuron input statistics for the test distribution may be different from the neuron input statistics for the train distribution, causing internal covariate shift, thus performance degradation as well as worsened uncertainty estimation.

In [8], the authors propose the test-time adaptation method Prediction-Time Batch Normalization (BN) to do inference on batches of test data and compute the neuron input statistics from the test batch. Instead of using the frozen estimates of neuron input statistics for the training data, the neuron input statistics for the test batch are used. A bigger test batch size will naturally result in better estimates of the neuron input statistics for the true test distribution. BN as defined by [8] does not have any hyperparameters.

Test entropy minimization (TENT). A visual recognition model for n classes running inference on test image X , outputs a vector Y representing probabilities y_i for each possible class i . The entropy of this prediction:

$$H(Y) = - \sum_{i=1}^n P(y_i) \log P(y_i) \quad (3)$$

is a measure for the confidence of the model on its prediction Y for test image X . A lower entropy corresponds to a more confident prediction. In [16], the authors find that test entropy correlates with distribution shift, where more entropy correlates with a bigger shift. Based on this finding, they propose the test-time adaptation method Test Entropy Minimization (TENT).

TENT builds on top of BN; inference is done on batches of test data and the neuron input statistics for the test batch are used instead of those estimated during training. The objective in TENT is to minimize test entropy on a batch of test images, by updating the transformation parameters in the batch normalization layers of the model. More specifically, the parameters γ and β for all batch normalization layers are updated using gradient descent, on every newly encountered test batch. Unlike the neuron input statistics, which are computed while forwarding the test batch through the network, the transformation parameters are updated after the forward pass. The only hyperparameter of TENT is its learning rate.

TENT may suffer from error accumulation. When the model makes an incorrect prediction, TENT will learn from the erroneous prediction. In the next batches, this may lead to the model making more of such errors. As TENT is unsupervised, it cannot recover from this error reinforcement.

Continual test-time adaptation (CoTTA). In a continually changing test distribution setting, error accumulation may become a significant problem. The authors of [17] propose Continual Test-Time Adaptation (CoTTA) to tackle the issues of previous works on test-time adaptation on a continually changing target. CoTTA builds on top of TENT. Instead of only updating the transformation parameters, CoTTA updates all parameters in the model and does not only rely on batch normalization layers for its adaptation. The effects of error accumulation are reduced by updating the model parameters with an exponential moving average:

$$\theta'_{t+1} = \alpha * \theta'_t + (1 - \alpha) * \theta_{t+1} \quad (4)$$

where a higher value for the smoothing factor α discounts older batches faster, resulting in faster adaptation, but more risk of error accumulation.

Furthermore, with a continually changing target, TENT suffers from catastrophic forgetting. After adapting to the continually changing target, the model may not perform as well anymore on the source distribution. To circumvent this problem, CoTTA stochastically restores source model weights. The restoration hyperparameter p determines the probability that a weight is restored to its original value from the source model. A higher value of p decreases forgetting, but limits the model from adapting.

4 Experimental Setup

We evaluate test-time adaptation methods in the continual test-time adaptation setting under contextual and semantic shifts, for which we use CORE50 [7] and DomainNet [10] respectively. For both datasets, all images have a label for their class as well as for their domain. Different splits of source and target domains result in different experiments. For each experiment, the source domains are further split into a 90% train and 10% validation (hold-out) set. A source model is trained on the train set of the source domains. Whenever we do evaluation on source domain(s) (e.g. to measure forgetting), we do this on the 10% validation sets only. Whenever we do evaluation on target domain(s), we do this on all the data available for those domains.

For test-time adaptation method CoTTA, we always run it with the default learning rate and disable test-time augmentations.

4.1 Contextual Shifts on CORE50

Dataset. For evaluation of test-time adaptation methods under continual contextual shifts, we perform experiments on the CORE50 dataset [7]. The dataset consists of 50 domestic objects as classes. These objects are captured in 11 distinct contextually different domains with different backgrounds and lighting conditions. For each object and domain, there are 300 images (frames) from a 15 seconds video of the object. The images are from the point-of-view of the operator moving and rotating the objects in his hand slowly. The holding hand is consistent within a domain and is different between domains. The hand causes the object to be occluded in different ways. The authors of CORE50 did not design the dataset specifically for continual test-time adaptation. Instead, it was designed for supervised continual learning. They propose a split of domains (called sessions in [7]) for training and testing. Three of the eleven domains (3, 7 and 8) have been selected for testing and the remaining eight domains have been selected for training. The authors tried to balance the difficulty of the train and test domains with respect to the holding hand, background and indoor/outdoor.

Source models. Similarly to the source models for DomainNet-126, we train the source models with three different seeds for each source domain (combination) and evaluate on all three models for each experiment and report the average. The model architecture consists of a feature encoder followed by a fully connected layer as classifier. The feature encoder is based on the ResNet-50 [3] backbone. The models were initialized with ImageNet-1K [2] weights and trained with batch size 128.

4.2 Semantic Shifts on DomainNet

Dataset. For evaluation of test-time adaptation methods under continual semantic shifts, we perform experiments on the DomainNet dataset [10]. The original dataset consists of 345 classes in six domains. The authors of [12] find that labels of certain classes and domains are very noisy. They propose a less noisy subset of 126 classes and four domains, which we will refer to as DomainNet-126 as they do in [1]. The domains in DomainNet-126 are real images (18,523 images), paintings (30,032 images), cliparts (18,523 images) and sketches (24,147 images). Although all classes are present in each of the four domains, they have a semantically different appearance.

Source models. We use the pre-trained models from [1]. The authors provide models trained with three different seeds for each source domain. We evaluate on all three models for each experiment and report the average. The model architecture consists of a feature encoder followed by a fully connected layer as classifier. The feature encoder is based on the ResNet-50 [3] backbone. The authors of [1] follow SHOT [6] by adding a 256-dimensional bottleneck after the backbone and apply WeightNorm [13] on the classifier. The models were initialized with ImageNet-1K [2] weights and trained with batch size 128.

5 Analysis

In this section, we will report on the results of the following experiments:

- 5.1 A qualitative inspection of the test-time adaptation methods.
- 5.2 An investigation on the effect of batch size on test-time batch normalization statistics adaptation.
- 5.3 An evaluation of short-term adaptation performance of the test-time adaptation methods.
- 5.4 An evaluation of long-term adaptation performance of the test-time adaptation methods.

5.1 Qualitative Inspection

In Figure 1 we visualize what effect the different methods for test-time adaptation have on the predictions of a model. We do this for 10 example objects from CORE50.

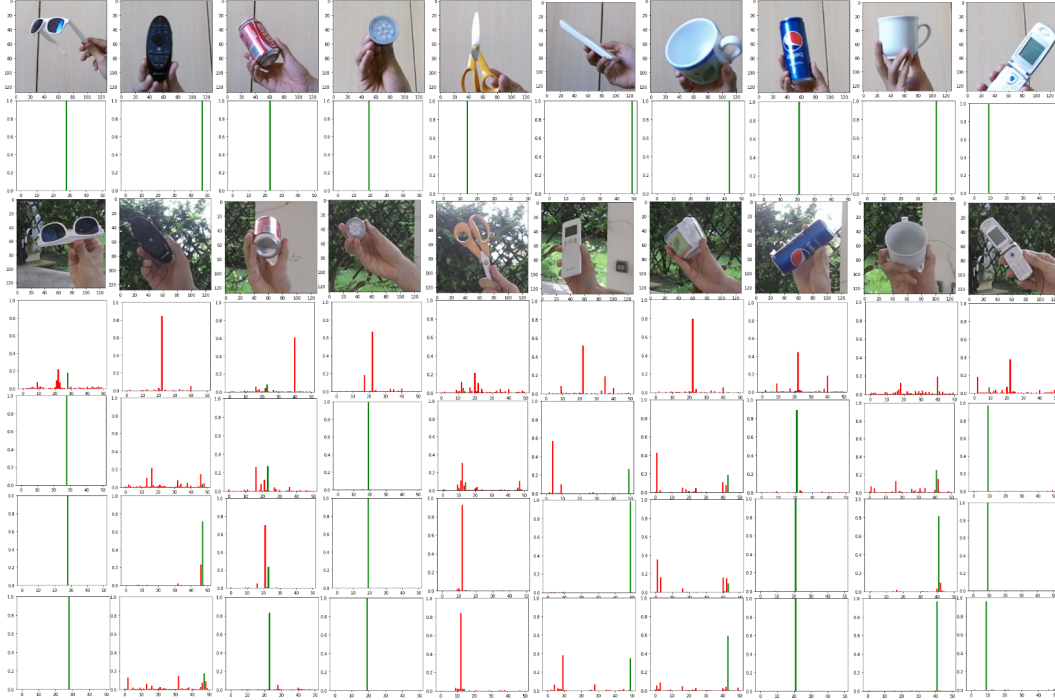


Figure 1: A qualitative inspection of the test-time adaptation methods on 10 example classes from CORE50. The first row shows images of 10 objects from indoor source domain 1. The second row shows the prediction confidence vectors of the source model (10/10 correct). The third row shows the same objects in outdoor target domain 10. The fourth row shows the prediction confidence vectors of the source model on the target domain (0/10 correct). The fifth row shows the prediction confidence vectors of BN on the target domain (6/10 correct). The sixth row shows the prediction confidence vectors of TENT on the target domain (7/10 correct). The last row shows the prediction confidence vectors of CoTTA on the target domain (8/10 correct).

In the top row, an example image is shown for each of the objects in the validation set of domain 1. Below that, in the second row, we visualize the prediction confidence vectors of a model trained only on domain 1 doing inference on these example images. The visualization is a bar chart, where the vertical axis represents prediction confidence and the horizontal axis represents all 50 possible classes. The bar corresponding to the correct class is colored green and all other bars are colored red. The final prediction of the model is simply the tallest bar, so the class with maximal confidence. The source model predicts all objects from the source domain validation set correctly with close to maximum confidence.

In the third row, an example image is shown for each of the objects in target domain 10. This domain is quite different from domain 1, as it is outdoor instead of indoor, resulting in contextual differences in background and lighting. Also the hand might hold the object differently, resulting in different occlusions of the objects. Below, in the fourth row, we see that the source model performs quite poorly due to this domain shift. It doesn't make one correct prediction on domain 10. On top of that, it is quite confident of its mistakes. So the uncertainty estimation of the model is not accurate either.

In the fifth row, we see the results of using BN. Simply updating the batch normalization statistics using a batch of test-time images leads to much more accurate predictions, of which 6 are correct now. The model also has better uncertainty estimation, the incorrect classes (red bars) have lower confidence (smaller bars). Please note that the effective batch size used is 400, so these ten images are part of a bigger batch of images from domain 10.

In the sixth row, we see the results of using TENT. Minimizing test entropy causes the model to make seven correct predictions, one more than BN. Please note that we divided domain 10 into batches of 400 images and that these ten example images are part of the last batch in the sequence. That means that TENT has been updating the batch normalization transformation parameters for many test batches in domain 10 before the predictions we see on the last batch. In general, we see that less classes get a significant confidence (less bars). The classes that do get significant confidence, are also higher in confidence. We see that some incorrect classes are predicted with high confidence, so TENT is making the uncertainty estimation less reliable. By minimizing entropy, TENT caused correct predictions to be more confident, but also accumulated errors by causing incorrect predictions to be more confident. We see TENT 'uncovering' the correct class in image two and six, but we also see TENT 'forgetting' the correct class in image three.

Finally, in the last row, we see the results of CoTTA, which updates all parameters of the network to minimize entropy, but uses weight averaging and stochastic restoration of the source model to reduce error accumulation and catastrophic forgetting. CoTTA causes the model to make eight correct predictions, one more than TENT. Compared to TENT, we see that CoTTA 'uncovered' the correct class in image three and seven, but failed or barely succeeded to 'uncover' the correct class in image two and six (which TENT did better). However, in these failure cases, we do see that the correct class is almost the same confidence as the predicted incorrect class and we hypothesize that after seeing more images from domain 10, these classes may get higher confidence over time.

5.2 Effect of Batch Size on Test-Time Batch Normalization Statistics Adaptation

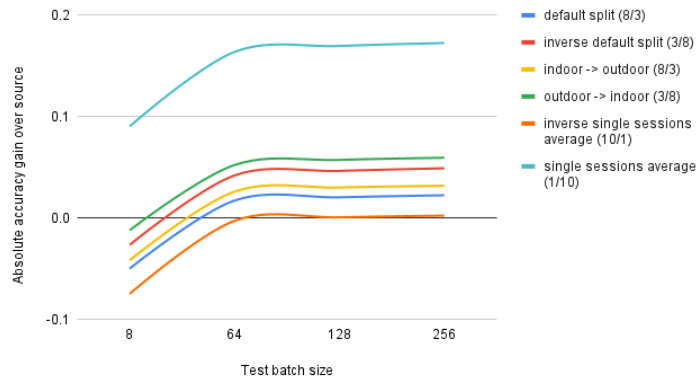


Figure 2: Effect of test batch size on classification accuracy (%) using the BN method for models trained on different source domains from CORE50.

All methods considered in our work rely on doing inference on batches of test images. Therefore, we investigate the effect of the test batch size on model accuracy when updating batch normalization statistics at test-time on CORE50 using the BN method. We evaluate models trained on different (combinations of) source domains on all 11 domains with and without BN for different batch sizes and we plot the average accuracy gain from BN over source in Figure 5.2. These results indicate that with more domains excluded from the source and included in the target, there is more accuracy gain. Furthermore, accuracy gain increases monotonically with batch size, with diminishing returns after

some point. In case of a small test batch size, accuracy may worsen using BN, due to inaccurate estimation of the statistics.

5.3 Short-Term Adaptation

With short-term adaptation, we refer to an experiment that will only see the same test-time image once. In the next section, we will go into long-term adaptation, where we repeat the experiment multiple times in a row. In this section, we evaluate short-term adaptation by first evaluating the source model on the validation set of the source domains (in-distribution), then sequentially evaluating it on the target domains (out-of-distribution) and finally evaluating it on the source domains again (in-distribution). We evaluate on the validation set before and after adapting to the target domains, such that we can measure forgetting. Forget rate is the difference in accuracy before and after the continual adaptation to the target domains.

5.3.1 Short-Term Adaptation to Contextual Shifts on CORE-50

Default train domains -> default test domains. In section 4 we explain what the default train/test split is in CORE50. If we train a source model on these train domains, the model should generalize quite well on the test domains. Despite the minimal distribution shifts, we still evaluate the test-time adaptation techniques on this split, as they should work well regardless of the presence and strength of shift. Furthermore, these results can be compared to other literature using the same split.

The results for adapting from the default train domains to the default test domains from CORE50 are shown in Table 3. Even though the test distribution was hypothesized to be relatively similar to the train distribution, BN still brings a significant improvement over the source model on sessions seven and ten. We see that TENT and CoTTA only bring small improvement over BN. We find that the small shift (and possibly short adaptation length) in this experiment will not lead to forgetting.

Table 3: Classification accuracy (%) for continual test-time adaptation on CORE50 [7] with source domains 1, 2, 4, 5, 6, 8, 9, 11.

Method	$t \rightarrow$					Mean	Forget rate
	1,2,4,5,6,8,9,11	3	7	10	1,2,4,5,6,8,9,11		
Source	100.0	83.6	54.0	59.0	100.0	79.3	0.0
BN [14]	100.0	83.6	68.4	66.9	100.0	83.8	0.0
TENT [16]	100.0	84.7	69.5	66.9	100.0	84.2	0.0
CoTTA [17]	100.0	83.7	69.5	67.3	100.0	84.1	0.0

Outdoor -> indoor. The results for continually adapting from outdoor to indoor domains from CORE50 are shown in Table 4. We find that for each target domain, BN leads to an improvement of accuracy. Furthermore, TENT leads to an even more significant improvement on all target domains. However, TENT also results in forgetting, causing a small accuracy drop on the source domain after the adaptation sequence. Interestingly, although CoTTA does consistently outperform BN, it never outperforms TENT for any of the domains. We hypothesize CoTTA needs more iterations to improve on TENT, but it is impossible to see that from this experiment.

Table 4: Classification accuracy (%) for short-term adaptation on CORE50 [7] with domains 4, 10, 11 as source and effective batch size of 400.

Method	$t \rightarrow$										Mean	Forget rate
	4,10,11	1	2	3	5	6	7	8	9	4,10,11		
Source	100.0	59.1	51.8	57.9	51.5	26.8	31.3	56.8	49.1	100.0	58.4	0.0
BN	100.0	64.1	60.1	59.8	62.9	59.1	43.1	64.3	53.7	100.0	66.7	0.0
TENT [16] (lr=8e-4)	100.0	68.8	68.0	71.5	73.3	67.5	55.8	73.9	71.3	98.1	74.8	1.9
CoTTA [17] ($\alpha=.99$, $p=.01$)	100.0	64.9	62.4	63.1	68.2	62.8	49.2	69.3	60.1	99.9	70.1	0.1

Single domain -> all other domains. The results for continually adapting from a single domain to all other domains from CORE50 are shown in Table 5. The findings are pretty similar to adapting from outdoor to indoor, however CoTTA does outperform TENT on some domains in this experiment. Another difference is that we see less forgetting in TENT and more in CoTTA, but nothing very significant.

Table 5: Classification accuracy (%) for short-term adaptation on CORE50 [7] with domain 1 as source (effective batch size is 400).

Method	$t \longrightarrow$											Mean	Forget rate	
	1	2	3	4	5	6	7	8	9	10	11			
Source	100.0	47.0	51.8	20.9	49.7	30.7	35.3	64.6	51.0	13.8	23.9	100.0	49.1	0.0
BN [14]	100.0	62.8	64.3	53.6	60.2	51.9	48.4	74.7	61.9	45.0	52.0	100.0	64.6	0.0
TENT [16] ($\text{lr}=3\text{e-}4$)	100.0	65.0	67.4	56.4	63.1	58.9	55.7	76.1	65.8	52.9	58.5	99.6	68.3	0.4
CoTTA [17] ($\alpha=.99$, $p=.01$)	100.0	64.5	66.3	57.8	63.2	57.5	54.0	74.7	66.1	52.0	59.2	99.2	67.9	0.8

5.3.2 Short-Term Adaptation to Semantic Shifts on DomainNet

Real images -> paintings, cliparts and sketches. The results for continually adapting from real images to paintings, cliparts and sketches from DomainNet-126 are shown in Table 6. We find that overall, BN performs comparable with source (with a very slight improvement overall), however for clipart specifically we see a small decrease in performance. TENT and CoTTA show only a small improvement over source. Noteworthy here is that CoTTA is showing double the forgetting. Updating all parameters in the model for CoTTA increases the effect of forgetting as opposed to only updating the batch normalization transformation parameters for TENT.

Table 6: Classification accuracy (%) for short-term adaptation on DomainNet [10] with real images as source (effective batch size is 400).

Method	$t \longrightarrow$					Mean	Forget rate
	<i>Real</i>	<i>Painting</i>	<i>Clipart</i>	<i>Sketch</i>	<i>Real</i>		
Source	98.2	62.7	55.5	46.4	98.2	72.2	0.0
BN [14]	98.2	63.4	54.3	47.7	98.2	72.4	0.0
TENT [16] ($\text{lr}=1\text{e-}4$)	98.2	65.1	57.0	53.2	95.9	73.9	2.4
CoTTA [17] ($\alpha=.99$, $p=.1$)	97.8	65.0	57.6	53.3	93.6	73.5	4.2

Paintings -> real images, cliparts and sketches. The results for continually adapting from paintings to real images, cliparts and sketches from DomainNet-126 are shown in Table 7. BN leads to a bigger improvement over source than for adapting from real images, albeit still quite small. TENT does only very slightly outperform BN, but shows some forgetting. CoTTA outperforms TENT on all target domains, however, it shows even more forgetting. We hypothesize that the shift from paintings to the other domains is bigger than when shifting from real images to the other domains, as real images show much more detail and will probably have more variety. Paintings can be seen as reconstructions of the real world, therefore almost never being as rich in information as a real image. Another possible reason is that the target domains together account for 54% more images when source is paintings instead of real images, making the adaptation longer, which could very well cause more forgetting as well.

Table 7: Classification accuracy (%) for short-term adaptation on DomainNet [10] with paintings as source (effective batch size is 400).

Method	$t \longrightarrow$					Mean	Forget rate
	<i>Painting</i>	<i>Real</i>	<i>Clipart</i>	<i>Sketch</i>	<i>Painting</i>		
Source	97.5	75.0	53.0	47.3	97.5	74.1	0.0
BN [14]	97.4	75.2	54.4	52.7	97.4	75.4	0.0
TENT [16] ($\text{lr}=1\text{e-}4$)	97.4	76.1	56.9	56.8	92.0	75.8	5.4
CoTTA [17] ($\alpha=.99$, $p=.1$)	97.3	77.8	59.4	57.8	84.0	75.3	13.4

5.4 Long-Term Adaptation

With long-term adaptation, we refer to an experiment that will see the same test-time image multiple times. In the previous section, we covered short-term adaptation, where we fix the length of the experiment by the number of unique test images. In this section, once again we start by first evaluating the source model on the validation set of the source domains (in-distribution), then sequentially evaluating it on the target domains (out-of-distribution) and finally evaluating it on the source domains again (in-distribution). Different from short-term adaptation however, to measure long-term adaptation, we repeat this cycle multiple times and refer to one cycle as an epoch. Even though we measure accuracy on the source data after every epoch, we first make a copy of the model under adaptation and use that to evaluate on the source data. This copy is thrown away after every epoch, so the actual model under adaptation never sees the source data again during the experiment. Also, much like in standard supervised training, we randomize the order of images in each target domain after every epoch, leading to different batches every epoch. We do not change the order of the target domains though, which is the same order as in the short-term adaptation experiments.

The results for long-term adaptation from real images to paintings, cliparts and sketches from DomainNet-126 are shown in Figure 3. We find that TENT quickly deteriorates in performance, whereas CoTTA keeps on improving. Also, TENT starts to forget more than linearly to the number of epochs, whereas for CoTTA forgetting is increased linearly and with relatively little amounts. We also see that CoTTA starts outperforming the maximum accuracy TENT can reach. Similar results have been found for the other source/target splits from section 5.3.

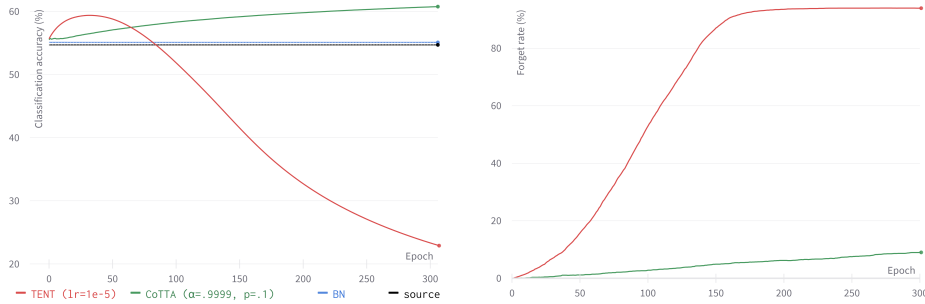


Figure 3: Long-term adaptation on DomainNet [10] with real images as source (effective batch size is 600).

6 Discussion

This section will go into some very important points of discussion related to our work.

Once a model is deployed, it may be doing inference for an infinite amount of time. Unless we know beforehand on how much images a model will run inference on, we cannot derive much meaningful conclusions from experiments with a fixed length.

TENT is very sensitive to its learning rate. We had to carefully tune it to the specific experiment and its length to work well. CoTTA works well on the default learning rate, however also for CoTTA the other hyperparameters had to be tuned to the specific experiment and its length to work well. This may not be possible in a real life scenario, where you do not know what type of shifts will happen or how long your model will be running inference for.

The short-term-adaptation experiments can be quite deceiving, leading one to believe that TENT is better than CoTTA. However, it seems CoTTA is better at long-term adaptation and just needs more iterations to start working. Therefore, in some of the short-term experiments, we set the alpha hyperparameter quite aggressively, causing CoTTA to show a lot of forgetting, whereas it would not have such extreme forgetting in a long-term setting with a lower alpha.

From the limited long-term-adaptation experiments, we cannot conclude whether CoTTA will maintain its high accuracy forever. This requires further work, as described in 8. It could be that by increasing the alpha hyperparameter, we are just postponing the error accumulation causing

the performance to decrease. That would mean we are still tuning it to the specific length of the experiment and there is no way to flatten the performance curve somehow.

Also, we did not try lower learning rates of TENT for the long-term adaptation experiments. In order to be sure that TENT cannot perform as well as CoTTA in the long-term, we therefore suggest in 8 to do more experiments on it.

In our experiments, we do not reset TENT (or CoTTA) whenever the target domain changes. In [17], the authors show that this will improve performance significantly. However, to stay closer to a real life scenario where information about domain changes is unavailable, we opt not to reset the model.

We turned off augmentations in CoTTA because we find that the augmentations provided by the authors significantly worsen performance on our experiments. We hypothesize this could be because the source models weren't trained with such augmentations either. Experiments we did on other datasets show that using the same augmentations used during training for CoTTA does lead to performance improvement.

7 Conclusion

In general, we find that the test-time adaptation methods show relatively more improvements on contextual shifts compared to semantic shifts. Also we find that the test-time-adaptation methods show more forgetting in semantic shifts compared to contextual shifts.

For relatively short fixed experiments lengths and carefully tuned hyperparameters for the specific experiment, we find that TENT overall outperforms the source model, BN and CoTTA.

From the (limited) long-term adaptation experiments, we find that CoTTA starts to outperform TENT when the experiment length is longer. These results indicate that CoTTA is better suited to a real-life scenario where a model is running inference for an infinite amount of time. We also find that the bigger the shift or the longer the experiment, the more forgetting occurs for both TENT and CoTTA. Forgetting seems to be linearly related to experiment length.

Finally, we can conclude that, BN, given a large enough test batch size, will perform on par with or outperform the source model (with only one small exception in our experiments). On top of that, BN does not show error accumulation or catastrophic forgetting, as it is not optimization based and will produce the same results on the same test batch irrespective of the previous test batches. The gains of BN seem to be lower when the source model was trained on a dataset with a relatively wide distribution.

8 Future Work

There are quite some areas of further work that are very important to pursue, in order to derive more sound and insightful conclusions. This section will go into some areas of future work in order of importance.

1. For the long-term adaptation experiments, different versions of TENT should be added to the graph with different learning rates, in order to conclude whether TENT cannot perform as well as CoTTA in the long-term.
2. The long-term adaptation experiments should be extended to even more epochs to find out whether CoTTA flattens or starts decreasing after a while.
3. Long-term adaptation versions should be added for the other experiments in the short-term experiments section.
4. Add experiments with other splits of source and target domains, such that we can base our conclusions on more results.
5. Do the batch size analysis also on DomainNet, as we saw examples where BN decreases performance which can almost not be explained by the very high effective batch size of 400. If there is another reason, find out what that is.
6. Do the qualitative analysis also on DomainNet.

7. Add experiments where the domain is shifting gradually instead of very abrupt distinct changes.
8. Turn on augmentations for CoTTA and analyze when which augmentations work and how this is related to the augmentations used for source model training. Augmentations may further improve CoTTA’s long-term adaptation performance, as every epoch CoTTA will see different augmentations of the same image, much like normal supervised training.

References

- [1] D. Chen, D. Wang, T. Darrell, and S. Ebrahimi. Contrastive test-time adaptation. *CVPR*, 2022.
- [2] J. Deng, W. Dong, R. Socher, L.-J. Li, K. Li, and L. Fei-Fei. Imagenet: A large-scale hierarchical image database. *CVPR*, 2009.
- [3] K. He, X. Zhang, S. Ren, and J. Sun. Deep residual learning for image recognition. *CVPR*, 2016.
- [4] S. Ioffe and C. Szegedy. Batch normalization: Accelerating deep network training by reducing internal covariate shift. *ICML*, 2015.
- [5] A. Khurana, S. Paul, P. Rai, S. Biswas, and G. Aggarwal. Sita: Single image test-time adaptation. 2021.
- [6] J. Liang, D. Hu, and J. Feng. Do we really need to access the source data? source hypothesis transfer for unsupervised domain adaptation. *ICML*, 2020.
- [7] V. Lomonaco and D. Maltoni. Core50: a new dataset and benchmark for continuous object recognition. *CoRL*, 2017.
- [8] Z. Nado, S. Padhy, D. Sculley, A. D’Amour, B. Lakshminarayanan, and J. Snoek. Evaluating prediction-time batch normalization for robustness under covariate shift. *ICML UDL*, 2020.
- [9] S. Niu, J. Wu, Y. Zhang, Y. Chen, S. Zheng, P. Zhao, and M. Tan. Efficient test-time model adaptation without forgetting. *ICML*, 2022.
- [10] X. Peng, Q. Bai, X. Xia, Z. Huang, K. Saenko, and B. Wang. Moment matching for multi-source domain adaptation. *ICCV*, 2019.
- [11] I. Redko, E. Morvant, A. Habrard, M. Sebban, and Y. Bennani. A survey on domain adaptation theory: learning bounds and theoretical guarantees. *CoRR2020*, 2020.
- [12] K. Saito, D. Kim, S. Sclaroff, T. Darrell, and K. Saenko. Semi-supervised domain adaptation via minimax entropy. *ICCV*, 2019.
- [13] T. Salimans and D. P. Kingma. Weight normalization: A simple reparameterization to accelerate training of deep neural networks. *NeurIPS*, 2016.
- [14] S. Schneider, E. Rusak, L. Eck, O. Bringmann, W. Brendel, and M. Bethge. Improving robustness against common corruptions by covariate shift adaptation. *NeurIPS*, 2020.
- [15] Y. Sun, X. Wang, Z. Liu, J. Miller, A. Efros, and M. Hardt. Test-time training with self-supervision for generalization under distribution shifts. *ICML*, 2020.
- [16] D. Wang, E. Shelhamer, S. Liu, B. Olshausen, and T. Darrell. Tent: Fully test-time adaptation by entropy minimization. *ICLR*, 2021.
- [17] Q. Wang, O. Fink, L. Van Gool, and D. Dai. Continual test-time domain adaptation. *CVPR*, 2022.
- [18] F. You, J. Li, and Z. Zhao. Test-time batch statistics calibration for covariate shift. 2021.
- [19] Y. Zou, Z. Zhang, C.-L. Li, H. Zhang, T. Pfister, and J.-B. Huang. Learning instance-specific adaptation for cross-domain segmentation. 2022.