

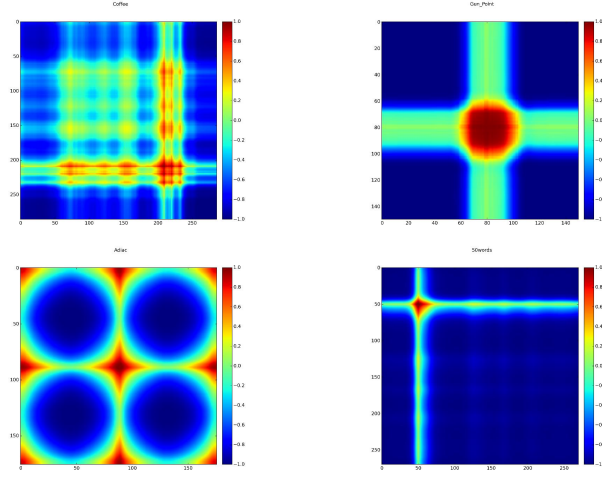


FIG. 4.2. Examples of GAF images on the 'Coffee', 'Gun-Point', 'Adiac' and '50Words' datasets.

the time axis. $w_{i,j}$ is given by the frequency with which a point in the quantile q_j is followed by a point in the quantile q_i . After normalization by $\sum_j w_{ij} = 1$, W is the Markov transition matrix. It is insensitive to the distribution of X and temporal dependency on time steps t_i . However, getting rid of the temporal dependency results in too much information loss in matrix W . To overcome this drawback, we define the Markov Transition Field (MTF) as follows:

$$(4.8) \quad M = \begin{bmatrix} w_{ij|x_1 \in q_i, x_1 \in q_j} & \cdots & w_{ij|x_1 \in q_i, x_n \in q_j} \\ w_{ij|x_2 \in q_i, x_1 \in q_j} & \cdots & w_{ij|x_2 \in q_i, x_n \in q_j} \\ \vdots & \ddots & \vdots \\ w_{ij|x_n \in q_i, x_1 \in q_j} & \cdots & w_{ij|x_n \in q_i, x_n \in q_j} \end{bmatrix}$$

We build a $Q \times Q$ Markov transition matrix (say, W) by dividing the data (magnitude) into Q quantile bins. The quantile bins that contain the data at time stamp i and j (temporal axis) are q_i and q_j ($q \in [1, Q]$). M_{ij} in MTF denotes the transition probability of $q_i \rightarrow q_j$. That is, we spread out matrix W which contains the transition probability on magnitude

axis into the MTF matrix by considering the temporal positions.

By assigning the probability from the quantile at time step i to the quantile at time step j at each pixel M_{ij} , the MTF M actually encodes the multi-span transition probabilities of the time series. $M_{i,j||i-j|=k}$ denotes the transition probability between the points with time interval k . For example, $M_{ij|j-i=1}$ illustrates the transition process along the time axis with a skip step. The main diagonal M_{ii} , which is a special case when $k = 0$ captures the probability from each quantile to itself (the self-transition probability) at time step i . To make the image size manageable and computation more efficient, we reduce the MTF size by averaging the pixels in each non-overlapping $m \times m$ patch with the blurring kernel $\{\frac{1}{m^2}\}_{m \times m}$. That is, we aggregate the transition probabilities in each subsequence of length m together. Figure 4.3 shows the procedure to encode time series to MTF.

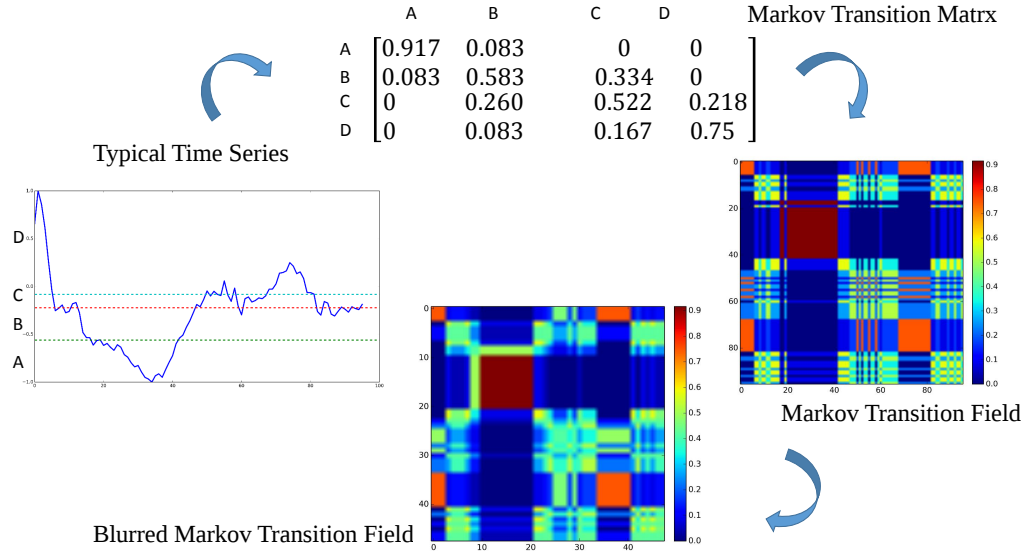


FIG. 4.3. Illustration of the proposed encoding map of Markov Transition Field. X is a sequence of typical time series in dataset 'ECG'. X is firstly discretized into Q quantile bins. Then we calculate its Markov Transition Matrix W and finally build its MTF with eq. (4.8). In addition, we reduce the image size from 96×96 to 48×48 by averaging the pixels in each non-overlapping 2×2 patch.

More full MTF images are shown in Figure 4.4.

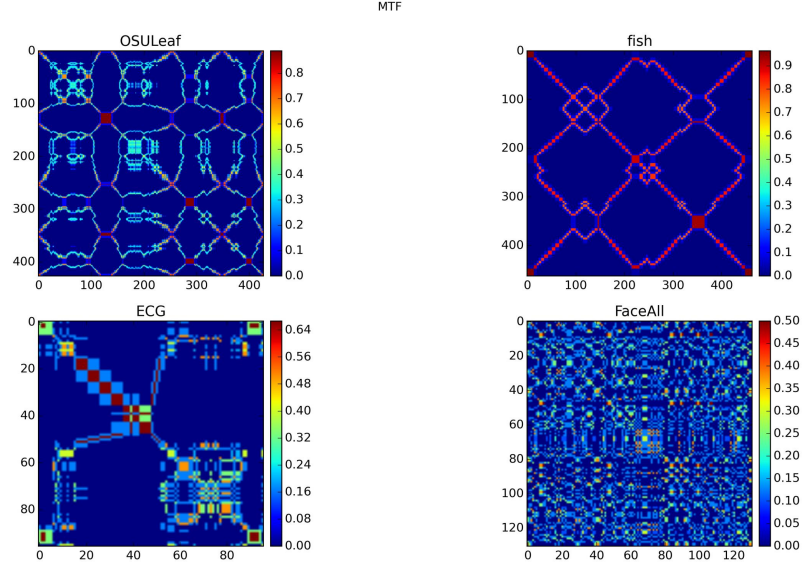


FIG. 4.4. Example of full MTF on 4 dataset (without blurring) with 16 quantile bins.

4.5 Classifying Time Series Using Deep Learning Architectures

In this section, we will explore the classification performance on time series images through using different deep learning (DL) architectures which include Convolutional Neural Networks (CNN), Tiled Convolutional Neural Networks (TCNN), Stacked Autoencoders (SAE) and Deep Belief Nets (DBN). It is also interesting to compare the performance and learned features of these DL methods on our non-natural images.

4.5.1 Tiled Convolutional Neural Networks

Tiled Convolutional Neural Networks (Ngiam *et al.* 2010) are a variation of Convolutional Neural Networks. They use tiles and multiple feature maps to learn invariant features. Tiles are parameterized by a tile size K to control the distance over which weights

are shared. By producing multiple feature maps, Tiled CNNs learn overcomplete representations through unsupervised pretraining with Topographic ICA (TICA).

A typical TICA network is actually a double-stage optimization procedure with square and square root nonlinearities in each stage, respectively. In the first stage, the weight matrix W is learned while the matrix V is hard-coded to represent the topographic structure of units. More precisely, given a sequence of inputs $\{x^{(h)}\}$, the activation of each unit in the second stage is $f_i(x^{(h)}; W, V) = \sqrt{\sum_{k=1}^p V_{ik} (\sum_{j=1}^q W_{kj} x_j^{(h)})^2}$. TICA learns the weight matrix W in the second stage by solving:

$$(4.9) \quad \begin{aligned} & \underset{W}{\text{minimize}} \quad \sum_{h=1}^n \sum_{i=1}^p f_i(x^{(h)}; W, V) \\ & \text{subject to} \quad WW^T = I \end{aligned}$$

$W \in \mathbb{R}^{p \times q}$ and $V \in \mathbb{R}^{p \times p}$ where p is the number of hidden units in a layer and q is the size of the input. V is a logical matrix ($V_{ij} = 1$ or 0) that encodes the topographic structure of the hidden units by a contiguous 3×3 block. The orthogonality constraint $WW^T = I$ provides diversity among learned features.

The pretraining algorithm (Algorithm. 1) is based on gradient descent on the TICA objective function in Equation. 4.9. The inner loop is a simple implementation of backtracking linesearch. The *orthogonalize_localRF*(W^{new}) function only orthogonalizes the weights that have completely overlapping receptive fields. Weight-tying is applied by averaging each set of tied weights. The algorithm is trained by batch projected gradient descent. Other unsupervised feature learning algorithms such as RBMs and autoencoders (Bengio *et al.* 2007) require more parameter tuning, especially during optimization. However, pretraining with TICA usually requires little tuning of optimization parameters, because the tractable objective function of TICA allows to monitor convergence easily.

Algorithm 1 Unsupervised pretraining with TICA.

Require: $\{x^{(t)}\}_{t=1}^T, v, s, W, V$ as input

Ensure: W as output

repeat

$$f^{old} = \sum_{t=1}^T \sum_{i=1}^m \sqrt{\sum_{k=1}^m V_{ik} (\sum_{j=1}^n W_{kj} x_j^{(t)})^2}, g = \frac{f^{old}}{\partial W}, f^{new} = +\infty, \alpha = 1$$

while $f^{new} > f^{old}$ **do**

$$W^{new} = W - \alpha g$$

$$W^{new} = Localize(W^{new}, s)$$

$$W^{new} = tieWeights(W^{new}, k)$$

$$W^{new} = orthogonalizeLocalRF(W^{new})$$

$$W^{new} = tieWeights(W^{new}, k)$$

$$f^{new} = \sum_{t=1}^T \sum_{i=1}^m \sqrt{\sum_{k=1}^m V_{ik} (\sum_{j=1}^n W_{kj} x_j^{(t)})^2}$$

$$\alpha = 0.5\alpha$$

end while

$$W = W^{new}$$

until convergence

Neither GAF nor MTF images are natural images; they have no natural concepts such as “edges” and “angles”. Thus, we propose to exploit the benefits of unsupervised pretraining with TICA to learn many diverse features with local orthogonality. In (Ngiam *et al.* 2010), the authors empirically demonstrate that tiled CNNs perform well with limited labeled data because the partial weight tying requires fewer parameters and reduces the need for a large amount of labeled data. Our data from the UCR Time Series Repository (Keogh *et al.* 2011) tends to have few instances (e.g., the “yoga” dataset has 300 labeled instance in the training set and 3000 unlabeled instance in the test set), so tiled CNNs are suitable for our learning task. Moreover, Tiled CNNs achieve good performance on large datasets (such as NORB and CIFAR). W

Typically, tiled CNNs are trained with two hyperparameters, the tiling size k and the number of feature maps l . In our experiments, we directly fixed the network structures without tuning these hyperparameters in loops, since the configuration is proved to be the

optimal in their work. Our experimental settings follow the default deep network structures and parameters in (Ngiam *et al.* 2010). Tiled CNNs with such configurations are reported to achieve the best performance on the NORB image classification benchmark. Although tuning the parameters will surely enhance performance, doing so may cloud our understanding of the power of the representation. Another consideration is computational efficiency. All of the experiments on the 12 datasets could be done in one day on a laptop with an Intel i7-3630QM CPU and 8GB of memory (our experimental platform). Thus, the results in this paper are a preliminary lower bound on the potential best performance. Thoroughly exploring network structures and parameters will be addressed in future work. The structure and parameters of the tiled CNN used in this paper are illustrated in Figure 4.5.

4.5.2 Experimental Settings and Data

We apply Tiled CNNs to classify time series using GASF, GADF and MTF representation on 20 datasets from UCR time series repository. More detailed statistics are summarized in Table 4.1. The datasets are pre-split into training and testing sets for experimental comparisons. For each dataset, the table gives its name, the number of classes, the number of training and test instances, and the length of the individual time series.

We apply Tiled CNNs to classify time series using GAF and MTF representations on 20 datasets from (Keogh *et al.* 2011) in different domains such as medicine, entomology, engineering, astronomy, signal processing, and others. The datasets are pre-split into training and testing sets to facilitate experimental comparisons. We compare the classification error rate of our GASF-GADF-MTF approach with previously published results of 6 best approaches proposed recently: Fast-Shapelets(Rakthanmanon & Keogh 2013), a 1NN classifier based on SAX with Bag-of-Patterns (SAX-BoP) (Lin, Khade, & Li 2012), a SAX based Vector Space Model (SAX-VSM)(Senin & Malinchik 2013), a classifier based on the

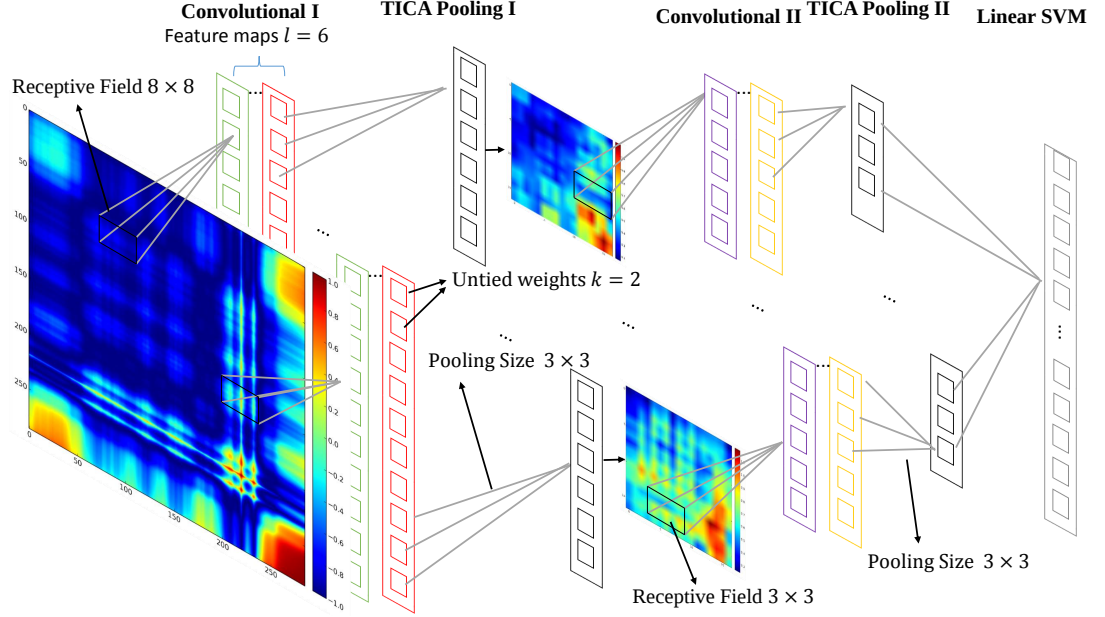


FIG. 4.5. Structure of the tiled convolutional neural networks. We fix the size of receptive fields to 8×8 in the first convolutional layer and 3×3 in the second convolutional layer. Each TICA pooling layer pools over a block of 3×3 input units in the previous layer without warping around the borders to optimize for sparsity of the pooling units. The number of pooling units in each map is exactly the same as the number of input units. The last layer is a linear SVM for classification. We construct this network by stacking two Tiled CNNs, each with 6 maps ($l = 6$) and tiling size $k = 1, 2, 3$.

Recurrence Patterns Compression Distance (RPCD) (Silva *et al.* 2013), a tree-based symbolic representation for multivariate time series (SMTS) (Baydogan & Runger 2014) and a SVM classifier based on a bag-of-features representation (TSBF) (Baydogan, Runger, & Tuv 2013).

In our experiments, the size of GAF image is regulated by the the number of PAA bins S_{GAF} . Given a time series X of size n , we divide the time series into S_{GAF} adjacent, non-overlapping windows along the time axis and extract the means of each bin. This enables us to construct the smaller GAF matrix $G_{S_{GAF} \times S_{GAF}}$. MTF requires the time series to be discretized into Q quantile bins to calculate the $Q \times Q$ Markov transition matrix, from which

Table 4.1. Statistics on 20 UCR datasets

	classes	train	test	length
50Words	50	450	455	270
Adiac	37	390	391	176
Beef	5	30	30	470
CBF	3	30	900	128
Coffee	2	28	28	286
ECG	2	100	100	96
Face(all)	14	560	1690	131
Face(four)	4	24	88	350
Fish	7	175	175	463
Gun-Point	2	50	150	150
Lighting-2	2	60	61	637
Lighting-7	7	70	73	319
OliveOil	4	30	30	570
OSUL.	6	200	242	427
SwedishL.	15	500	625	128
Synt.Cont.	6	300	300	60
Trace	4	100	100	275
TwoPat.	4	1000	4000	128
Wafer	2	1000	6174	152
Yoga	2	300	3000	426

we construct the raw MTF image $M_{n \times n}$ afterwards. Before classification, we shrink the MTF image size to $S_{MTF} \times S_{MTF}$ by the blurring kernel $\{\frac{1}{m^2}\}_{m \times m}$ where $m = \lceil \frac{n}{S_{MTF}} \rceil$. The Tiled CNN is trained with image size $\{S_{GAF}, S_{MTF}\} \in \{16, 24, 32, 40, 48\}$ and quantile size $Q \in \{8, 16, 32, 64\}$. At the last layer of the Tiled CNN, we use a linear soft margin SVM (Fan *et al.* 2008) and select C by 5-fold cross validation over $\{10^{-4}, 10^{-3}, \dots, 10^4\}$ on the training set.

For each input of image size S_{GAF} or S_{MTF} and quantile size Q , we pretrain the Tiled CNN with the full unlabeled dataset (both training and test set) to learn the initial weights W through TICA. Then we train the SVM at the last layer by selecting the penalty factor C with cross validation. Finally, we classify the test set using the optimal hyperparam-

Table 4.2. Summary of error rates for 3 classic baselines, 6 recently published best results and our approach. The symbols $\triangleleft$, $*$, $\dagger$ and $\bullet$ represent datasets generated from human motions, figure shapes, synthetically predefined procedures and all remaining temporal signals, respectively. For our approach, the numbers in brackets are the optimal PAA size and quantile size.

Dataset	Random-Guess	Fast-Shapelet	SAX-BoP	SAX-VSM	RPCD	SMTS	TSBF	GASF-GADF-MTF
50words $\bullet$	0.98	N/A	0.466	N/A	0.2264	0.289	0.209	0.301 (16, 32)
Adiac $*$	0.972	0.514	0.432	0.381	0.3836	0.248	0.245	0.373 (32, 48)
Beef $\bullet$	0.8	0.447	0.433	0.33	0.3667	0.26	0.287	0.233 (64, 40)
CBF $\dagger$	0.667	0.053	0.013	0.02	N/A	0.02	0.009	0.009 (32, 24)
Coffee $\bullet$	0.5	0.068	0.036	0	0	0.029	0.004	0 (64, 48)
ECG $\bullet$	0.5	0.227	0.15	0.14	0.14	0.159	0.145	0.09 (8, 32)
FaceAll $*$	0.75	0.411	0.219	0.207	0.1905	0.191	0.234	0.237 (8, 48)
FaceFour $*$	0.98	0.090	0.023	0	0.0568	0.165	0.051	0.068 (8, 16)
fish $*$	0.857	0.197	0.074	0.017	0.1257	0.147	0.08	0.114 (8, 40)
Gun.Point $\triangleleft$	0.5	0.061	0.027	0.007	0	0.011	0.011	0.08 (32, 32)
Lighting2 $\bullet$	0.5	0.295	0.164	0.196	0.2459	0.269	0.257	0.114 (16, 40)
Lighting7 $\bullet$	0.857	0.403	0.466	0.301	0.3562	0.255	0.262	0.260 (16, 48)
OliveOil $\bullet$	0.75	0.213	0.133	0.1	0.1667	0.177	0.09	0.2 (8, 48)
OSULeaf $*$	0.833	0.359	0.256	0.107	0.3554	0.377	0.329	0.358 (16, 32)
SwedishLeaf $*$	0.933	0.269	0.198	0.01	0.0976	0.08	0.075	0.065 (16, 48)
synthetic control $\dagger$	0.8338	0.081	0.037	0.251	N/A	0.025	0.008	0.007 (64, 48)
Trace $\dagger$	0.75	0.002	0	0	N/A	0	0.02	0 (64, 48)
Two Patterns $\dagger$	0.75	0.113	0.129	0.004	N/A	0.003	0.001	0.091 (64, 32)
wafer $\bullet$	0.5	0.004	0.003	0.0006	0.0034	0	0.004	0 (64, 16)
yoga $*$	0.5	0.249	0.17	0.164	0.134	0.094	0.149	0.196 (8, 32)
#wins	0	0	1	5	3	4	4	9

ters $\{S, Q, C\}$ with the lowest error rate on the training set. If two or more models tie, we prefer the larger S and Q because larger S helps preserve more information through the PAA procedure and larger Q encodes the dynamic transition statistics with more detail. Our model selection approach provides generalization without being overly expensive computationally.

4.5.3 Results and Discussion

We use Tiled CNNs to classify the single GASF, GADF and MTF images as well as the compound GASF-GADF-MTF images on 20 datasets. For the sake of space, we do not show the full results on single-channel images. Generally, our approach is not prone to overfitting by the relatively small difference between training and test set errors. One exception is the Olive Oil dataset with the MTF approach where the test error is significantly higher.

In addition to the risk of potential overfitting, we found that MTF has generally higher error rates than GAFs. This is most likely because of the uncertainty in the inverse map of MTF. Note that the encoding function from $-1/1$ rescaled time series to GAFs and MTF are both surjections. The map functions of GAFs and MTF will each produce only one image with fixed S and Q for each given time series X . Because they are both surjective mapping functions, the inverse image of both mapping functions is not fixed. However, the mapping function of GAFs on $0/1$ rescaled time series are bijective. As shown in a later section, we can reconstruct the raw time series from the diagonal of GASF, but it is very hard to even roughly recover the signal from MTF. Even for $-1/1$ rescaled data, the GAFs have smaller uncertainty in the inverse image of their mapping function because such randomness only comes from the ambiguity of $\cos(\phi)$ when $\phi \in [0, 2\pi]$. MTF, on the other hand, has a much larger inverse image space, which results in large variations when we try to recover the signal. Although MTF encodes the transition dynamics which are important features of time series, such features alone seem not to be sufficient for recognition/classification tasks.

Note that at each pixel, G_{ij} denotes the superstition/difference of the directions at t_i and t_j , M_{ij} is the transition probability from the quantile at t_i to the quantile at t_j . GAF encodes static information while MTF depicts information about dynamics. From this

point of view, we consider them as three “orthogonal” channels, like different colors in the RGB image space. Thus, we can combine GAFs and MTF images of the same size (i.e. $S_{GAFs} = S_{MTF}$) to construct a triple-channel image (GASF-GADF-MTF). It combines both the static and dynamic statistics embedded in the raw time series, and we posit that it will be able to enhance classification performance. In the experiments below, we pretrain and tune the Tiled CNN on the compound GASF-GADF-MTF images. Then, we report the classification error rate on test sets. In Table 4.2, the Tiled CNN classifiers on GASF-GADF-MTF images achieved significantly competitive results with 9 other state-of-the-art time series classification approaches.

4.6 Image Recovery on GASF for Time Series Imputation with Denoised Auto-encoder

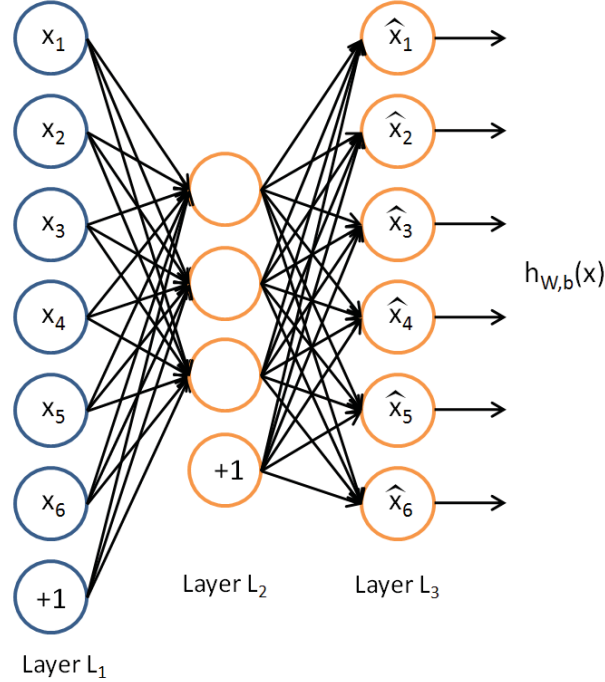


FIG. 4.6. Network structure of Auto-encoder.

We use Denoised AutoEncoders (DAE) to learn a generative model of time series on GASF. DAE is a variation of a standard AutoEncoder (Figure 4.6). An Autoencoder neural network is an unsupervised learning algorithm that applies backpropagation, setting the target values to be equal to the inputs.

The Auto-encoder tries to learn a function $h_{W,b}(x) \approx x$. In other words, it is trying to learn an approximation to the identity function, so as to output $\hat{x}$ that is similar to x . The identity function seems a particularly trivial function to be trying to learn; but by placing constraints on the network, such as by limiting the number of hidden units, we can discover interesting structure in the data. If the input were completely random say, each x_i comes from an IID Gaussian independent of the other features, then this compression task would be very difficult. But if there is structure in the data, for example, if some of the input features are correlated, then this algorithm will be able to discover these correlations (Ng 2011).

The idea behind Denoising Auto-encoders (DAE) is simple and similar to AE. In order to force the hidden layer to discover more robust features and prevent it from simply learning the identity, we train the Auto-encoder to reconstruct the input from a corrupted version of it. The DAE is a stochastic version of the auto-encoder. Intuitively, a DAE does two things: try to encode the input (preserve the information about the input), and try to undo the effect of a corruption process stochastically applied to the input of the AE. The latter can only be done by capturing the statistical dependencies between the inputs. The DAE can be understood from different perspectives including the manifold learning perspective, stochastic operator perspective, bottom-up information theoretic perspective and top-down generative model perspective (Vincent *et al.* 2008). The stochastic corruption process in this section randomly sets some of the inputs (as many as half of them) to zero (salt-and-pepper noise). Hence the DAE is trying to predict the corrupted (i.e., missing) values from the uncorrupted (i.e., non-missing) values, for randomly selected subsets of

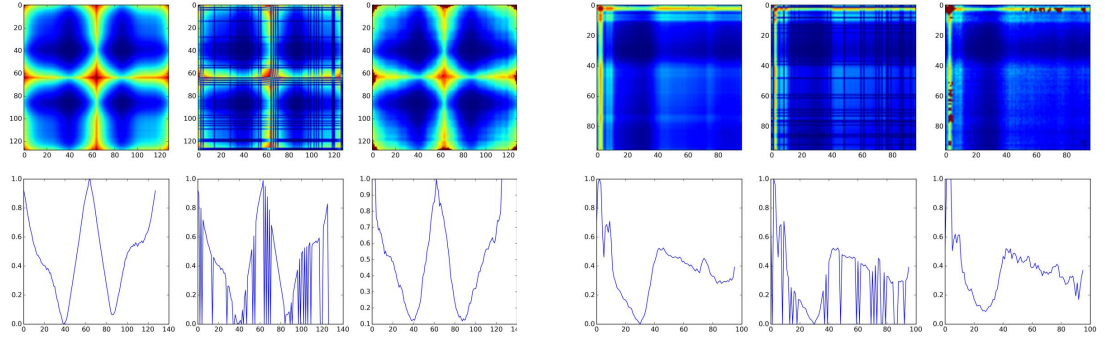


FIG. 4.7. Pipeline of time series imputation by image recovery. Raw GASF $\rightarrow$ "broken" GASF $\rightarrow$ recovered GASF (top), Raw time series $\rightarrow$ corrupted time series with missing value $\rightarrow$ predicted time series (bottom) on dataset "SwedishLeaf" (left) and "ECG" (right).

missing patterns. Note how being able to predict any subset of variables from the rest is a sufficient condition for completely capturing the joint distribution between a set of variables.

As previously mentioned, the mapping functions from $-1/1$ rescaled time series to GAFs are surjections. The uncertainty among the inverse images come from the ambiguity of the $\cos(\phi)$ when $\phi \in [0, 2\pi]$. However the mapping functions of $0/1$ rescaled time series are bijections. The main diagonal of GASF, i.e. $\{G_{ii}\} = \{\cos(2\phi_i)\}$ allows us to precisely reconstruct the original time series by

$$(4.10) \quad \cos(\phi) = \sqrt{\frac{\cos(2\phi) + 1}{2}} \quad \phi \in [0, \frac{\pi}{2}]$$

Thus, we can predict missing values among time series through recovering the "broken" GASF images. During training, we manually add "salt-and-pepper" noise (i.e., randomly set a number of points to 0) to the raw time series and transform the data to GASF images. Then a single layer Denoised Auto-encoder (DA) is fully trained as a generative model to reconstruct GASF images. Note that at the input layer, we do not add noise again to the "broken" GASF images. A Sigmoid function helps to learn the nonlinear features at

the hidden layer. At the last layer we compute the Mean Square Error (MSE) between the original and "broken" GASF images as the loss function to evaluate fitting performance. To train the models simple batch gradient descent is applied to back propagate the inference loss. For testing, after we corrupt the time series and transform the noisy data to "broken" GASF, the trained DA helps recover the image, on which we extract the main diagonal to reconstruct the recovered time series. To compare the imputation performance, we also test standard DA with the raw time series data as input to recover the missing values (Figure. 4.7).

4.6.1 Experiment Setting

For the DA models we use batch gradient descent with a batch size of 20. Optimization iterations run until the MSE changed less than a threshold of 10^{-3} for GASF and 10^{-5} for raw time series. A single hidden layer has 500 hidden neurons with sigmoid functions. We choose four dataset of different types from the UCR time series repository for the imputation task: "Gun Point" (human motion), "CBF" (synthetic data), "SwedishLeaf" (figure shapes) and "ECG" (other remaining temporal signals). To explore if the statistical dependency learned by the DA can be generalized to unknown data, we use the above four datasets and the "Adiac" dataset together to train the DA to impute two totally unknown test datasets, "Two Patterns" and "wafer" (We name these synthetic miscellaneous datasets "7 Misc"). To add randomness to the input of DA, we randomly set 20% of the raw data among a specific time series to be zero (salt-and-pepper noise). Our experiments for imputation are implemented with Theano (Bastien *et al.* 2012). To control for the random initialization of the parameters and the randomness induced by gradient descent, we repeated every experiment 10 times and report the average MSE.

Table 4.3. MSE of imputation on time series using raw data and GASF images.

Dataset	Full MSE		Interpolation MSE	
	Raw	GASF	Raw	GASF
ECG	0.01001	0.01148	0.02301	0.01196
CBF	0.02009	0.03520	0.04116	0.03119
Gun Point	0.00693	0.00894	0.01069	0.00841
SwedishLeaf	0.00606	0.00889	0.01117	0.00981
7 Misc	0.06134	0.10130	0.10998	0.07077

4.6.2 Results and Discussion

In Table 4.3, "Full MSE" means the MSE between the complete recovered and original sequence and "Imputation MSE" means the MSE of only the unknown points among each time series. Interestingly, DA on the raw data perform well on the whole sequence, generally, but there is a gap between the full MSE and imputation MSE. That is, DA on raw time series can fit the known data much better than predicting the unknown data (like overfitting). Predicting the missing value using GASF always achieves slightly higher full MSE but the imputation MSE is reduced by 12.18%-48.02%. We can observe that the difference between the full MSE and imputation MSE is much smaller on GASF than on the raw data. Interpolation with GASF has more stable performance than on the raw data.

Why does predicting missing values using GASF have more stable performance than using raw time series? Actually, the transformation maps of GAFs are generally equivalent to a kernel trick. By defining the inner product $k(x_i, x_j)$, we achieve data augmentation by increasing the dimensionality of the raw data. By preserving the temporal and spatial information in GASF images, the DA utilizes both temporal and spatial dependencies by considering the missing points as well as their relations to other data that has been explicitly encoded in the GASF images. Because the entire sequence, instead of a short subsequence, helps predict the missing value, the performance is more stable as the full MSE and imputation MSE.

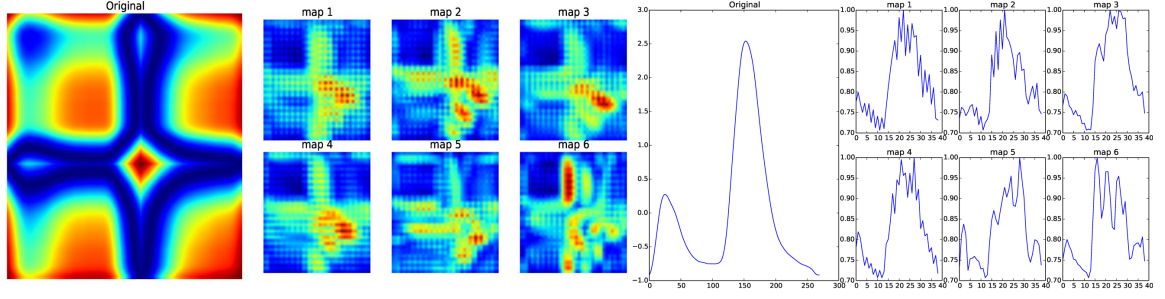


FIG. 4.8. (a) Original GASF and its six learned feature maps before the SVM layer in Tiled CNNs (left). (b) Raw time series and its reconstructions from the main diagonal of six feature maps on '50Words' dataset (right).

tation MSE are close.

4.7 Analysis on Features and Weights Learned by Tiled CNNs and DA

In contrast to the cases in which the CNNs is applied in natural image recognition tasks, neither GAFs nor MTF have natural interpretations of visual concepts like “edges” or “angles”. In this section we analyze the features and weights learned through Tiled CNNs to explain why our approach works.

Figure 4.8 illustrates the reconstruction results from six feature maps learned through the Tiled CNNs on GASF (by Eqn 4.10). The Tiled CNNs extracts the color patch, which is essentially a moving average that enhances several receptive fields within the nonlinear units by different trained weights. It is not a simple moving average but the synthetic integration by considering the 2D temporal dependencies among different time intervals, which is a benefit from the Gramian matrix structure that helps preserve the temporal information. By observing the orthogonal reconstruction from each layer of the feature maps, we can clearly observe that the tiled CNNs can extract the multi-frequency dependencies through the convolution and pooling architecture on the GAF and MTF images to preserve the trend while addressing more details in different subphases. The high-levelled feature maps

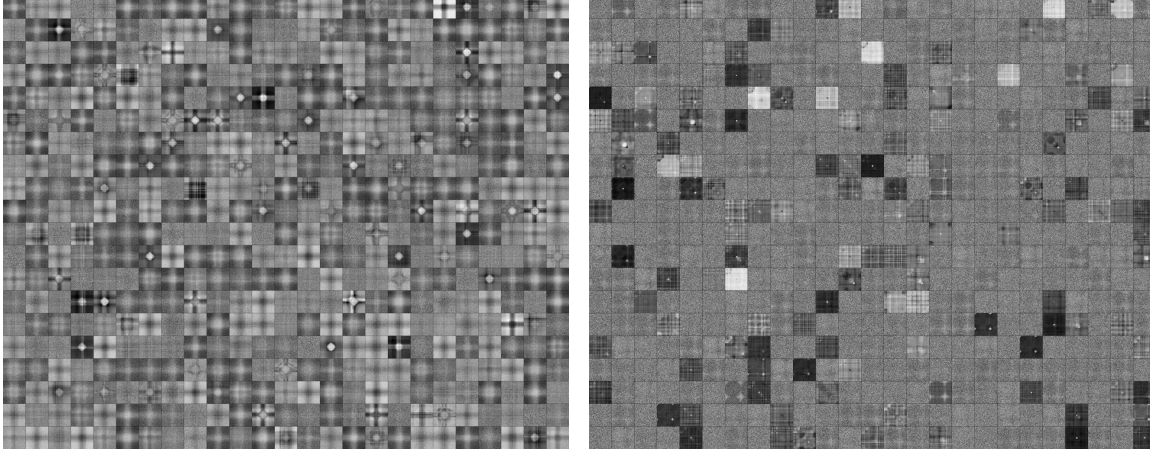


FIG. 4.9. All 500 filters learned by DA on the "Gun Point" (left) and "7 Misc" (right) dataset.

learned by the Tiled CNN are equivalent to a multi-frequency approximator of the original curve. Our experiments also demonstrates the learned weight matrix W with the constraint $WW^T = I$, which makes effective use of local orthogonality. The TICA pretraining provides the built-in advantage that the function w.r.t the parameter space is not likely to be ill-conditioned as $WW^T = 1$. The weight matrix W is quasi-orthogonal and approaching 0 without large magnitude. This implies that the condition number of W approaches 1 and helps the system to be well-conditioned.

As for imputation, because the GASF images have no concept of "angle" and "edge", DA actually learned different prototypes of the GASF images (Table 4.9). We find that there is significant noise in the filters on the "7 Misc" dataset because the training set is relatively small to better learn different filters. Actually, all the noisy filters with no patterns work like one Gaussian noise filter.

4.8 Experiments on Trajectory Data

We have demonstrated the effectiveness of GAF and MTF the benchmark time series datasets as diverse as shape, physiological surveillance and industry from the UCR time series repository. In this section we describe an application of our approaches to classify spatial-temporal trajectory data. The trajectory data is complex because patterns of movement are often driven by unperceived goals and constrained by an unknown environment.

To compare our results with other benchmark approaches including the seminal work from (Lee *et al.* 2008), we run experiments on two benchmark datasets, the animal movement dataset (Animal) and the hurricane track dataset (Hurricane) (Figure 4.10). Both datasets have trajectories of unequal length. For the "Animal" dataset, the x and y coordinates are extracted from animal movements observed in June 1995. It is divided into three classes by species: elk, deer, and cattle, as shown in Figure 15. The numbers of trajectories (points) are 38 (7117), 30 (4333), and 34 (3540), respectively. In the "Hurricane" dataset, the latitude and longitude are extracted from Atlantic hurricanes for the years 1950 through 2006. The Saffir-Simpson scale classifies hurricanes into categories 1-5 by intensity. A high category number indicates high intensity. Categories 2 and 3 are chosen for two classes. The numbers of trajectories (points) are 61 (2459) and 72 (3126), respectively. Both datasets are pre-split into two parts for training (80%) and testing (20%). Figure 4.10 shows the overview of the trajectory data. Table 4.4 provides the classes, training size, testing size, minimum length and maximum length of the trajectory data.

Table 4.4. Summary statistics of two trajectory datasets.

Dataset	Classes	Training Size	Testing Size	Min Length	Max Length
Animal Tracking	3	80	18	10	291
Hurricane	2	112	21	11	108

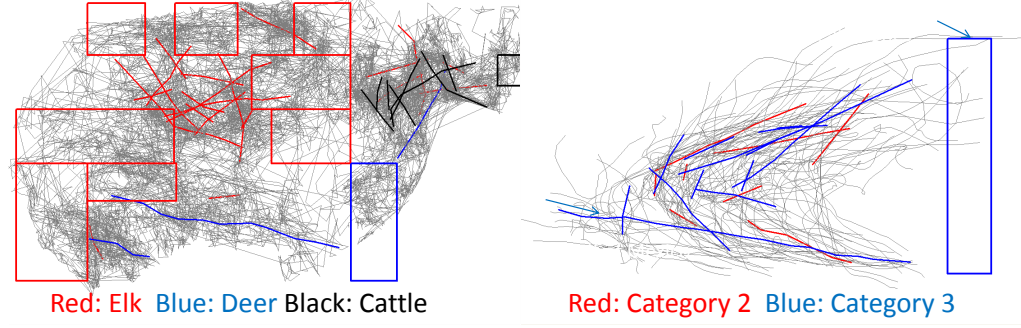


FIG. 4.10. Overview of the trajectory and the RB-TB features learnt in (a). Animal tracking data (left) and (b). Hurricane data (right)

4.8.1 Hilbert Space Filling Curves

Spatial-temporal trajectory data is commonly multi-dimensional. We use Hilbert Space Filling Curves (SFC) to transform the trajectory into time series while preserving the spatial-temporal information.

Space filling has been studied by the mathematicians since the late 19th century when the first graphical representation was proposed by David Hilbert in 1891 (Hilbert 1891). Space filling curves provide a linear mapping from the multi-dimensional space to the 1-dimensional space. This mapping can be thought of as dividing D-dimensional space into D-dimensional hypercubes with a line passing through each hypercube. Recently, filling curve based approaches have shown to be able to preserve locality between objects in the multidimensional space in the linear space, and thus have been applied to different tasks like clustering (Moon *et al.* 2001), high dimensional outlier detection (Angiulli & Pizzuti 2005), and trajectory query (Ding, Trajcevski, & Scheuermann 2008) and classification (Gandhi & Oates 2015). Figure 4.11 (a) shows SFC examples of order $\{1,2,3,4,5,6\}$.

Basically, the SFC of order 1 divides the square into 4 area. For the Hilbert curve with order 2, each sub-area of the curve with order 1 is further divided into 4 sub-areas. This

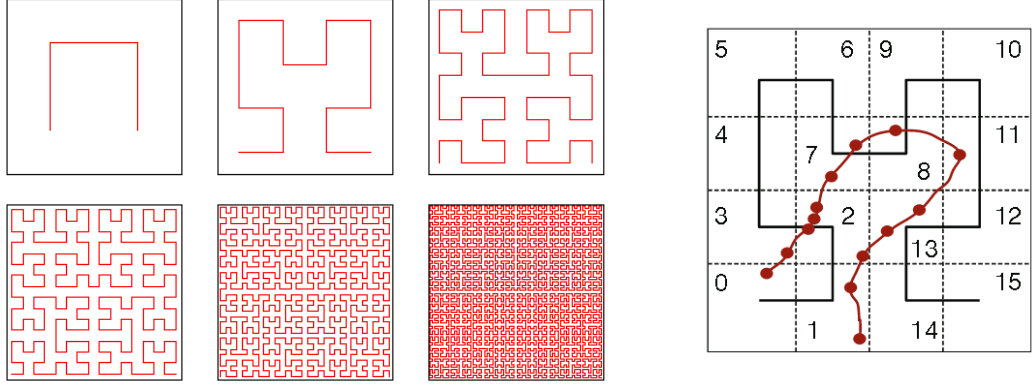


FIG. 4.11. (a). Hilbert space filling curve of order $\{1,2,3,4,5,6\}$ in 2-dimensional space (left) (b). An example of the transformation from 2-dimensional trajectory to 1-dimensional time series using HSCF of order 2 (right).

process goes on as the order of the SFC increases. It is clear that the number of sub-areas in 2 dimensional SFC is 4^{order} . To convert 2-dimensional data points to 1-dimensional points, each sub-area is integer numbered from 0 to $4^{order} - 1$ starting from the lower left corner as 0 to the lower right corner. All other sub-areas are numbered in order of occurrence of the corresponding vertex as shown in Figure 4.11 (b) when order = 2. It also shows the example transformation process from a 2D trajectory to a sequence of scalars (time series). The final time series generated after SFC transformation is $T = [0, 3, 2, 2, 2, 7, 7, 8, 11, 13, 13, 2, 1, 1]$.

We map the trajectory points by the visiting order of the SFC embedded in the trajectory manifold space to the index sequence by the recorded times. The produced time series can be used for classification using our algorithm. This adds another hyperparameter called the SFC order, which decides the granularity of the space filling curve.

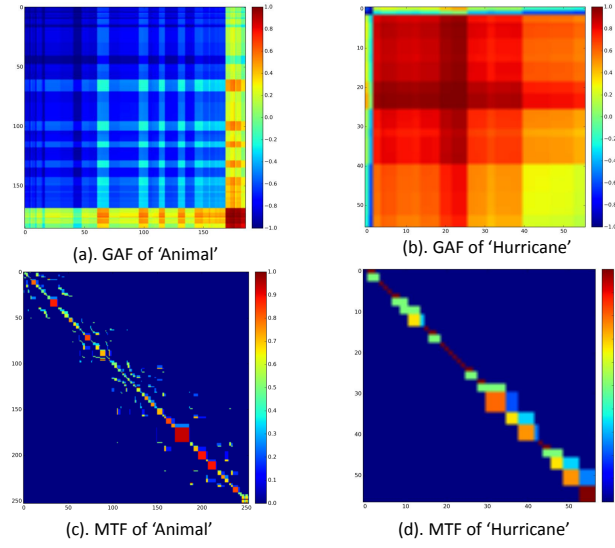


FIG. 4.12. Examples of GAF and MTF images generated from the time series on 'Animal' and 'Hurricane' datasets. The time series is produced using SFC from raw 2D trajectory.

4.8.2 Experiment Settings

The parameter settings are the same as the previous experiments on UCR datasets. The optimal SFC order is selected together with other parameters through 5-fold cross validation from $\{3,4,5,6,7,8,9,10\}$.

Note that both trajectory datasets have quite small sample size with varying length. When the trajectory length (as well as the time series length produced by SFC) is smaller than image size S , we uniformly duplicate each point in the time series in temporal order to stretch the sequence to length S . If the difference between the length of a time series and S is smaller than the original time series length, the interpolation strategy changes to random duplication instead of following the temporal order.

4.8.3 Results and Discussion

Both 'Animal' and 'Hurricane' datasets have been used in previous research (Lee *et al.* 2008; Gandhi & Oates 2015) to achieve state-of-the-art classification accuracy. Traclass give two algorithms, trajectory-based (TB-only) and region-based + trajectory-based (RB-TB) approaches based on features used for classification on these datasets. They carefully designed a hierarchy of features by partitioning trajectories and exploring two types of clustering. In (Gandhi & Oates 2015), the author used SFC transformation to linearly map the trajectory data to time series and classified the sequences based on symbolic discretization with the multiple normal distribution assumption.

After transforming the 2D trajectory data to time series using SFC, we generate the corresponding GAF and MTF images as shown in Figure 4.12. However, we found significant overfitting with CNNs even using 5-fold cross validation. This is probably because both the sample size and the time series length of the trajectory datasets are too small to avoid overfitting in neural networks. Previous work has discussed overfitting during cross validation and proposed potential techniques to address this problem (Ng 1997; Prechelt 1998). Here, we applied a simple and straight-forward hyperparameter selection approach to reduce classifier variance. For a given set of hyperparameter $\{S, Q, SFC_{order}\}$, after cross validation with different C values of the linear SVM, we compute the mean and standard deviation to get the 3σ lower bound over all C by

$$(4.11) \quad score_{3\sigma} = mean(Accuracy) - 3 \times STD(Accuracy)$$

By selecting the other hyperparameters $\{S, Q, SFC - order\}$ with the best statistical lower bound on the classifier performance over C , the optimal hyperparameters have lower variance while preserving lower bias. Using this hyperparameter selection approach, the

Table 4.5. Classification accuracy for TB-Only, RB-TB methods, multiple normal distribution based symbolic distance (NDist) and our algorithm (%).

Dataset	TB-Only	RB-TB	NDist	GAF-MTF
Animal Tracking	50	83.3	83.3	72.2
Hurricane	65.4	73.1	52.3	71.42

classification results are reported in Table 4.5.

We perform better than the TB-Only method on both datasets and almost as good as the RB-TB method on the 'Hurricane' dataset. However, both RB-TB and NDist methods outperform ours on the 'Animal' dataset. As shown in Figure 4.10, both region and trajectory based features are useful for classification. For the 'Hurricane' dataset, direction based features are more useful than region based features. Direction based features are quite easy to capture using our approach as the GAF is actually calculating the pairwise direction fields on each points in the trajectory data. For the 'Animal' dataset, region is very important as shown in Figure 4.10 (a). Elk, deer and cattle are almost separable just using location as their regions are clearly located at the left, right top and right bottom, respectively. When transforming the trajectory data into time series using SFC, two close regions might be mapped to different sub-areas with different SFC indexes. When the indexes of two close regions are also near, this can be handled by CNNs with its capability to capture the small shifting-invariance features. However, CNNs are not good at discriminating similar images with large shifting from each other. Thus, when the region information is preserved by the manner of shifting the specific patterns largely in the time series produced by SFC, CNNs might have difficulty capturing the region information.

Although our approach does not overtake other benchmark methods on both trajectory datasets, we provide a more general framework to encode the spatial-temporal patterns for classification tasks. Instead of complicated hand-tuned features, our approach can be

applied to a variety of time series and trajectory data. When the region of the trajectory is not significantly important or the direction feature dominates, our general methods work quite well. On large datasets where the volume of time series/trajectory data is big, our deep neural network based approach will greatly benefit from the large sample size in both feature learning and classification tasks.

4.9 Conclusion

This chapter described an off-line approach to spatially encode the temporal patterns for classification using convolutional neural networks. We created a pipeline for converting trajectory and time series data into novel representations, GAF and MTF images, and extracted high-level features from these using CNNs. The features were subsequently used for classification. We demonstrated that our approach yields competitive results when compared to state-of-the-art methods by searching a relatively small parameter space. We found that GAF-MTF multi-channel images are scalable to larger numbers of quasi-orthogonal features that yield more comprehensive images. Our analysis of high-level features learned from CNNs suggested Tiled CNNs work like multi-frequency moving averages that benefit from the 2D temporal dependency that is preserved by the Gramian matrix.

Chapter 5

SIGNIFICANT STATISTICS — ADAPTIVE NORMALIZED RISK-AVERTING TRAINING FOR NEURAL NETWORKS

5.1 Introduction

In last section, we investigated modeling temporal data through imaging its static and dynamic temporal correlations to learn the feature sets adaptively with the deep learning approaches. The significant non-linearity in DNNs make the training hard with a high chance to get stuck with local optima and plateaus. Meanwhile, the gradient vanishing/explosion problem exist when not using the Rectified Linear Unit, which can be alleviated by unsupervised pretraining techniques(Erhan *et al.* 2010). Some recent works point out that in high dimensional space (e.g. deep neural nets), the convexity is not needed because the local optima is good enough (Dauphin *et al.* 2014; Choromanska *et al.* 2014). However, all above techniques and works do not really demystify the non-convex optimization problem in DNNs and Recurrent Neural Networks (RNNs).

In this and next section, we will introduce a exponential-form based error estimator, whose optimality and robustness can be adjusted and achieved by tuning a index λ .

This new estimator provides another perspective to debunk the non-convexity in DNNs and RNNs.

5.2 Background

Deep neural networks (DNNs) are attracting attention largely due to their impressive empirical performance in image and speech recognition tasks. (Krizhevsky, Sutskever, & Hinton 2012; Hinton *et al.* 2012; Hannun *et al.* 2014). While Convolutional Networks (ConvNets) are the de facto state-of-the-art for visual recognition, Deep Belief Networks (DBN), Deep Boltzmann Machines (DBM) and Stacked Auto-encoders (SA) provide insights as generative models to learn the full generating distribution of input data. (Vincent *et al.* 2010; Hinton, Osindero, & Teh 2006; Salakhutdinov & Hinton 2009). Recently, researchers have investigated various techniques to improve the learning capacity of DNNs. Unsupervised pretraining using Restrict Boltzmann Machines (RBM), Denoised Autoencoders (DA) or Topographic ICA (TICA) has proved to be helpful for training DNNs with better weight initialization (Ngiam *et al.* 2010; Coates & Ng 2011). Specific types of deep network structures such as Network in Network (NIN) (Min Lin 2014) and Deeply Supervised Nets (DSN) (Lee *et al.* 2014) enhance the local and global modeling to enhance the feature learned by hidden layers. Rectified Linear Unit (ReLU) and variants are proposed as the optimal activation functions to better interpret hidden features (Nair & Hinton 2010; He *et al.* 2015). Various regularization techniques such as dropout (Srivastava *et al.* 2014) with Maxout (Goodfellow *et al.* 2013b) are proposed to regulate the DNNs to be less prone to overfitting.

Neural network models always lead to a non-convex optimization problem. The optimization algorithm impacts the quality of the local minimum because it is hard to find a global minimum or estimate how far a particular local minimum is from the best possible

solution. The most standard approach to optimize DNNs is Stochastic Gradient Descent (SGD). There are many variants of SGD and researchers and practitioners typically choose a particular variant empirically. While nearly all DNNs optimization algorithms in popular use are gradient-based, recent work has shown that more advanced second-order methods such as L-BFGS and Saddle-Free Newton (SFN) approaches can yield better results for DNN tasks (Ngiam *et al.* 2011; Dauphin *et al.* 2014). Second order derivatives can be addressed by hardware extensions (GPUs or clusters) or batch methods when dealing with massive data, SGD still provides a robust default choice for optimizing DNNs.

Instead of modifying the network structure or optimization techniques for DNNs, we focused on designing a new error function to convexify the error space. The convexification approach has been studied in the optimization community for decades, but has never been seriously applied within deep learning. Two well-known methods are the graduated non-convexity method (Blake & Zisserman 1987) and the LiuFloudas convexification method (Liu & Floudas 1993). LiuFloudas convexification can be applied to optimization problems where the error criterion is twice continuously differentiable, although determining the weight α of the added quadratic function for convexifying the error criterion involves significant computation when dealing with massive data and parameters.

Following the same name employed for deriving robust controllers and filters (Speyer, Deyst, & Jacobson 1974), a new type of Risk-Averting Error (RAE) is proposed theoretically for solving non-convex optimization problems (Lo 2010). Empirically, with the proposal of Normalized Risk-Averting Error (NRAE) and the Gradual Deconvexification method (GDC), this error criterion is proved to be competitive with the standard mean square error (MSE) in single layer and two-layer neural networks for solving data fitting and classification problems (Gui, Lo, & Peng 2014; Lo, Gui, & Peng 2012). Our work will consistently follow the paradigm of Lo's work. Interestingly, SimNets, a generalization of ConvNets that was recently proposed in (Cohen & Shashua 2014), uses the MEX

operator (whose name stands for Maximum-minimum-Expectation Collapsing Smooth) as an activation function to generalize ReLU activation and max pooling. We notice that the MEX operator with L_2 units has exactly the *same* mathematical form with NRAE. However, NRAE is still hard to optimize in practice due to plateaus and the unstable error space caused by the fixed large convexity index. GDC alleviates these problems but its performance is limited and suffers from the slow learning speed. Instead of fixing the convexity index λ , Adaptive Normalized Risk-Averting Training (ANRAT) optimizes NRAE by tuning λ adaptively using gradient descent. We give theoretical proofs of its optimal properties against the standard L_p -norm error. Our experiments on MNIST and CIFAR-10 with different deep/shallow neural nets demonstrate the effectiveness empirically. Being an optimization algorithm, our approach are not supposed to deal specifically with the problem of over-fitting, however we show that this can be handled by the usual methods of regularization such as weight decay or dropout.

In this section, we generalize RAE and NRAE to L_P norm and provide theoretical proofs on the convexity properties. We prove that NRAE is at least as good as MSE in solving Non-convex optimization problems when $|\lambda| \geq 1$. We develop an Adaptive Normalized Risk-Averting Training approach (ANRAT) to automatically tune the sensitive index λ through standard gradient descent. By analyzing the derivatives on λ , we show that our approach is analogous to Deconvexification by GDC while possessing more flexibility and fast convergence speed. Empirically, our approaches are able to overcome the under-fitting problem encountered when training DNN more effectively than the pretraining + fine-tuning. Being an optimization algorithm, our approach are not supposed to deal specifically with the problem of over-fitting, however we show that this can be handled by the usual methods of regularization such as weight decay or dropout. Our experiments demonstrate competitive results with the state-of-the-art on several visual recognition tasks. On MNIST dataset without pretraining, we can achieve 0.52% error rate using only ConvNets

(LeNet5) with weight decaym which is the best results ever reported with only standard ConvNets. Our preliminary experiments on CIFAR-10 dataset show that by using ConvNets and dropout with ANRAT methods, we can achieve 82.12% test accuracy, which is superior to the results using Cross Entropy (CE) and dropout (Zeiler & Fergus 2013) as well as competitive with the results achieved by unsupervised pretraining (Coates & Ng 2011). The experiments on Denoised Auto-encoder also demonstrate ANRAT helps in unsupervised feature learning.

5.3 Reformulation on Error Criterion - Significant Statistics

We begin with the definition of RAE for the L_p norm and the theoretical justifications on its convexity property. RAE is not suitable for real applications since it is not bounded. Instead, NRAE is bounded to overcome the register overflow in real implementations. We prove that NRAE is quasi-convex, and thus shares the same global and local optimum with RAE. Moreover, we show the lower-bound of its performance is as good as L_p -norm error when the convexity index satisfies a constraint, which theoretically supports the ANRAT method proposed in the next section.

5.3.1 Risk-averting Error Criterion

Given training samples $\{\mathbf{X}, y\} = \{(\mathbf{x}_1, y_1), (\mathbf{x}_2, y_2), \dots, (\mathbf{x}_m, y_m)\}$, $f(\mathbf{x}_i, \mathbf{W})$ is the learning model with parameters $\mathbf{W}$. The loss function of L_p -norm error is defined as:

$$(5.1) \quad l_p(f(\mathbf{x}_i, \mathbf{W}), y_i) = \frac{1}{m} \sum_{i=1}^m \|f(\mathbf{x}_i, \mathbf{W}) - y_i\|^p$$

When $p = 2$, Eqn. 6.1 denotes to the standard Mean Square Error (MSE). The Risk-

Averting Error criterion (RAE) corresponding to the L_p -norm error is defined by

$$(5.2) \quad RAE_{p,q}(f(\mathbf{x}_i, \mathbf{W}), y_i) = \frac{1}{m} \sum_{i=1}^m e^{\lambda^q \|f(\mathbf{x}_i, \mathbf{W}) - y_i\|^p}$$

λ is the convexity index. It controls the size of the convexity region.

For clarity, we write the matrix derivatives as function form. In our proof we use quadratic form.

The Jacobian matrix of Eqn. 6.1 is

$$(5.3) \quad J_{p,q}(\mathbf{W}) = \frac{1}{m} p \lambda^q \sum_{i=1}^m e^{\lambda^q \|f(\mathbf{x}_i, \mathbf{W}) - y_i\|^p} \|f(\mathbf{x}_i, \mathbf{W}) - y_i\|^{p-1} \frac{\partial f(\mathbf{x}_i, \mathbf{W})}{\partial \mathbf{W}}$$

The Hessian matrix of Eqn. 6.1 is

$$(5.4) \quad H_{p,q}(\mathbf{W}) = \frac{p}{m} \sum_{i=1}^m e^{\lambda^q (f(\mathbf{x}_i, \mathbf{W}) - y_i)^p}$$

$$(5.5) \quad \times \{ (p-1) \lambda^q \|f(\mathbf{x}_i, \mathbf{W}) - y_i\|^{p-2} \frac{\partial f(\mathbf{x}_i, \mathbf{W})}{\partial \mathbf{W}}^2$$

$$(5.6) \quad + p \lambda^{2q} (y_i - f(\mathbf{x}_i, \mathbf{W}))^{2p-2} \frac{\partial f(\mathbf{x}_i, \mathbf{W})}{\partial \mathbf{W}}^2$$

$$(5.7) \quad + \lambda^q (y_i - f(\mathbf{x}_i, \mathbf{W}))^{p-1} \frac{\partial f^2(\mathbf{x}_i, \mathbf{W})}{\partial \mathbf{W}^2} \}$$

Because RAE has the sum-exponential form, its Hessian matrix is tuned exactly by the convexity index λ^q . The following theorem indicates the relation between the convexity index and its convexity region.

Theorem 1 (Convexity). *Given the Risk-Averting Error criterion $RAE_{p,q}$ ($p, q \in \mathcal{N}^+$), which is twice continuous differentiable. $J_{p,q}(\mathbf{W})$ and $H_{p,q}(\mathbf{W})$ are the corresponding Jacobian and Hessian matrix. As $\lambda \rightarrow \pm\infty$, the convexity region monotonically expands to the entire parameter space except for the subregion $S := \{W \in \mathcal{R}^n | \text{rank}(H_{p,q}(\mathbf{W})) < n, H_{p,q}(\mathbf{W}) < 0\}$.*

Proof. Assume $p > 0$ and $q > 0$, matrix 5.4 is positive semi-definite. Note that both $\frac{\partial f(\mathbf{x}_i, \mathbf{W})^2}{\partial \mathbf{W}}$ and $\lambda^{2q}(y_i - f(\mathbf{x}_i, \mathbf{W}))^{2p-2}$ are positive semi-definite, matrix 6.7 is semi-positive definite, but Eqn. 6.6 and 6.8 may be indefinite. Let $\alpha_i(\mathbf{W}) = f(\mathbf{x}_i, \mathbf{W}) - y_i$, We rewrite Eqn. 6.7 in quadratic form:

$$(5.8) \quad = p\lambda^{2q}\alpha_i(\mathbf{W})^T \frac{\partial f(\mathbf{x}_i, \mathbf{W})^T}{\partial \mathbf{W}} \cdot \frac{\partial f(\mathbf{x}_i, \mathbf{W})}{\partial \mathbf{W}} \alpha_i(\mathbf{W})$$

$$(5.9) \quad = p\lambda^{2q}\alpha_i(\mathbf{W})^T Q \Lambda Q^T \alpha_i(\mathbf{W})$$

$$(5.10) \quad = p\lambda^{2q}S(\mathbf{W})^T \Lambda S(\mathbf{W})$$

$\Lambda = \text{diag}[\Lambda_1, \Lambda_2, \dots, \Lambda_m]$. Note that when p is a even number, Eqn. 6.6 is also positive semi-definite.

If $S(\mathbf{W})$ is a full-rank matrix, then Eqn. 6.9 is positive definite. When $\lambda^q \rightarrow \pm\infty$, the eigenvalues Λ becomes dominant in the leading principal minors (as well as eigenvalues) of the Hessian matrix to make $H_{p,q}(\mathbf{W})$ monotonically increasing with λ^q . Assume $P_\lambda := \{W \in \mathcal{R}^n | H_{p,q}(\mathbf{W}) > 0\}$, we have $P_{\lambda_1} \subset P_{\lambda_2}$ when $\lambda_1 < \lambda_2$. Considering Eqn. 6.6, 6.7 and 6.8 are all bounded, $\exists \psi(W)$, s.t. $H_{p,q}(\mathbf{W}) > 0$ when $|\lambda^q| > \psi(W)$.

When $S(\mathbf{W})$ is not a full-rank matrix, the determinant of all its $\binom{n}{k}$ submatrices is 0. Thus, in the subregion $S := \{W \in \mathcal{R}^n | \text{rank}(H_{p,q}(\mathbf{W})) < n, H_{p,q}(\mathbf{W}) < 0\}$, there is no parameters satisfied the convexity conditions ($\bigcup P_\lambda$). $\square$

Please refer to the supplementary material for the proof. Intuitively, the use of the RAE was motivated by its emphasizing large individual deviations in approximating functions and optimizing parameters in an exponential manner, thereby avoiding such large individual deviations and achieving robust performances. Theoretically, Theorem 7 states that when the convexity index λ increases to infinity, the convexity region in the parameter space of RAE expands monotonically to the entire space except the intersection of a finite

number of lower dimensional sets. The number of sets increases rapidly as the number m of training samples increases. Roughly speaking, larger λ and m cause the size of the convexity region to grow larger respectively in the error space of RAE.

When $\lambda \rightarrow \infty$, the error space can be perfectly stretched to be strictly convex, thus avoid the local optimum to guarantee a global optimum. Although RAE works well in theory, it is not bounded and suffers from the exponential magnitude and arithmetic overflow when using gradient descent in implementations .

5.3.2 Normalized Risk-Averting Error Criterion

RAE ensures the convexity of the error space to find the global optimum. By using NRAE, we relax the global optimum problem by finding a better local optimum to meet a theoretically and practically reasonable trade-off in real applications.

Given training samples $\{\mathbf{X}, y\} = \{(\mathbf{x}_1, y_1), (\mathbf{x}_2, y_2), \dots, (\mathbf{x}_m, y_m)\}$, $f(\mathbf{x}_i, \mathbf{W})$ is the learning model with parameters $\mathbf{W}$. The Normalized Risk-Averting Error Criterion (NRAE) corresponding to the L_p -norm error is defined as:

$$\begin{aligned}
 & NRAE_{p,q}(f(\mathbf{x}_i, \mathbf{W}), y_i) \\
 &= \frac{1}{\lambda^q} \log RAE_{p,q}(f(\mathbf{x}_i, \mathbf{W}), y_i) \\
 (5.11) \quad &= \frac{1}{\lambda^q} \log \frac{1}{m} \sum_{i=1}^m e^{\lambda^q \|f(\mathbf{x}_i, \mathbf{W}) - y_i\|^p}
 \end{aligned}$$

Theorem 2 (Bounded). $NRAE_{p,q}(f(\mathbf{x}_i, \mathbf{W}), y_i)$ is bounded. $NRAE_{p,q}(f(\mathbf{x}_i, \mathbf{W}), y_i) \rightarrow \minimax$ if $\lambda^q \rightarrow \infty$, $NRAE_{p,q}(f(\mathbf{x}_i, \mathbf{W}), y_i) \rightarrow L_p\text{-norm error}$ if $\lambda^q \rightarrow 0$.

Proof. Let

$$(5.12) \quad \alpha_i(\mathbf{W}) = f(\mathbf{x}_i, \mathbf{W}) - y_i$$

$$(5.13) \quad \alpha_{max}(\mathbf{W}) = \max\{\alpha_i(\mathbf{W}), i = 1, \dots, m\}$$

$$(5.14) \quad \beta_i(\mathbf{W}) = e^{\lambda^q(|\alpha_{max}(\mathbf{W})|^p - |\alpha_i(\mathbf{W})|^p)}$$

Then we can write Eqn. 5.11 as

$$\begin{aligned} NRAE_{p,q}(f(\mathbf{x}_i, \mathbf{W}), y_i) &= \frac{1}{\lambda^q} \log \frac{1}{m} e^{\lambda^q |\alpha_{max}|^p} \sum_{i=1}^m \beta_i(\mathbf{W}) \\ (5.15) \quad &= -\frac{1}{\lambda^q} \log m + |\alpha_{max}|^p + \frac{1}{\lambda^q} \log \sum_{i=1}^m \beta_i(\mathbf{W}) \\ &\leq -\frac{1}{\lambda^q} \log m + |\alpha_{max}|^p + \frac{1}{\lambda^q} \log m \cdot e^{\lambda^q |\alpha_{max}|^p} \\ (5.16) \quad &= 2|\alpha_{max}|^p \end{aligned}$$

□

The proof is provided in the supplemental materials. Briefly, NRAE is bounded by functions independent of λ and no overflow occurs for $\lambda \gg 1$. The following theorem states the quasi-convexity of NRAE.

Theorem 3 (Quasi-convexity). *Given a parameter space $\{W \in \mathcal{R}^n\}$, Assume $\exists \psi(W)$, s.t. $H_{p,q}(\mathbf{W}) > 0$ when $|\lambda^q| > \psi(W)$ to guarantee the convexity of $RAE_{p,q}(f(\mathbf{x}_i, \mathbf{W}), y_i)$. Then, $NRAE_{p,q}(f(\mathbf{x}_i, \mathbf{W}), y_i)$ is **quasi-convex** and share the same local and global optimum with $RAE_{p,q}(f(\mathbf{x}_i, \mathbf{W}), y_i)$.*

Proof. If $RAE_{p,q}(f(\mathbf{x}_i, \mathbf{W}), y_i)$ is convex, it is quasi-convex. $\log$ function is monotonically increasing, so the composition $\log RAE_{p,q}(f(\mathbf{x}_i, \mathbf{W}), y_i)$ is quasi-convex.¹

¹Because the function f defined by $f(x) = g(U(x))$ is quasi-convex if the function U is quasiconvex and

$\log$ is a strictly monotone function and $NRAE_{p,q}(f(\mathbf{x}_i, \mathbf{W}), y_i)$ is quasi-convex, so it shares the same local and global minimizer with $RAE_{p,q}(f(\mathbf{x}_i, \mathbf{W}), y_i)$. $\square$

The convexity region of NRAE is consistent with RAE. To interpret this statement in another perspective, the $\log$ function is a strictly monotone function. Even if RAE is not strictly convex, NRAE still shares the same local and global optimum with RAE. If we define the mapping function $f : RAE \rightarrow NRAE$, it is easy to see that f is bijective and continuous. Its inverse map f^{-1} is also continuous, so that f is an open mapping. Thus, it is easy to prove that the mapping function f is a homeomorphism to preserve all the topological properties of the given space.

The above theorems state the consistent relations among NRAE, RAE and MSE. It is proven that the greater the convexity index λ , the larger is the convex region is. Intuitively, increasing λ creates tunnels for a local-search minimization procedure to travel through to a good local optimum. However, we care about the justification on the advantage of NRAE against MSE. Theorem 9 provides the theoretical justification for the performance lower-bound of NRAE.

Theorem 4 (Lower-bound). *Given training samples $\{\mathbf{X}, y\} = \{(\mathbf{x}_1, y_1), (\mathbf{x}_2, y_2), \dots, (\mathbf{x}_m, y_m)\}$ and the model $f(\mathbf{x}_i, \mathbf{W})$ with parameters $\mathbf{W}$. If $\lambda^q \geq 1$, $p, q \in \mathcal{N}^+$ and $p \geq 2$, then both $RAE_{p,q}(f(\mathbf{x}_i, \mathbf{W}), y_i)$ and $NRAE_{p,q}(f(\mathbf{x}_i, \mathbf{W}), y_i)$ always have the higher chance to find a better local optimum than the standard L_p -norm error due to the expansion of the convexity region.*

Proof. Let $h_p(\mathbf{W})$ denotes the Hessian matrix of standard L_p -norm error (Eqn. 6.1), note

the function g is increasing.

$\alpha_i(\mathbf{W}) = f(\mathbf{x}_i, \mathbf{W}) - y_i$ we have

$$\begin{aligned}
 h_p(\mathbf{W}) &= \frac{p}{m} \sum_{i=1}^m \left\{ (p-1) \alpha_i(\mathbf{W})^{p-2} \frac{\partial f(\mathbf{x}_i, \mathbf{W})^2}{\partial \mathbf{W}} \right. \\
 (5.17) \quad &\quad \left. + \alpha_i(\mathbf{W})^{p-1} \frac{\partial f^2(\mathbf{x}_i, \mathbf{W})}{\partial \mathbf{W}^2} \right\}
 \end{aligned}$$

Since $\lambda^q \geq 1$, let $diag_{eig}$ denotes the diagonal matrix of the eigenvalues from SVD decomposition. $\succeq$ here means 'element-wise greater'. When $A \succeq B$, each element in A is greater than B. Then we have

$$\begin{aligned}
 diag_{eig}[H_{p,q}(\mathbf{W})] &\succeq diag_{eig}[h_p(\mathbf{W}) + \\
 &\quad \frac{p^2}{m} \sum_{i=1}^m \|\alpha_i(\mathbf{W})\|^{2p-2} \frac{\partial f(\mathbf{x}_i, \mathbf{W})^2}{\partial \mathbf{W}} \} \\
 (5.18) \quad &\succeq diag_{eig}[h_p(\mathbf{W})]
 \end{aligned}$$

This indicates that the $RAE_{p,q}(f(\mathbf{x}_i, \mathbf{W}), y_i)$ always has larger convexity regions than the standard L_p -norm error to better enable escape of local minima. Because $NRAE_{p,q}(f(\mathbf{x}_i, \mathbf{W}), y_i)$ is quasi-convex, sharing the same local and global optimum with $RAE_{p,q}(f(\mathbf{x}_i, \mathbf{W}), y_i)$, the above conclusions are still valid. $\square$

Roughly speaking, NRAE always has a larger convexity region than the standard L_p -norm error in terms of their Hessian matrix when $\lambda \geq 1$. This property guarantees the higher probability to escape poor local optima using NRAE. In the worst case, NRAE will perform as good as standard L_p -norm error if the convexity region shrinks as λ decreases or the local search deviates from the "tunnel" of convex regions.

More specifically, $NRAE_{p,q}(f(\mathbf{x}_i, \mathbf{W}), y_i)$

- approaches the standard L_p -norm error as $\lambda^q \rightarrow 0$.

- approaches the minimax error criterion $\inf_W \alpha_{max}(\mathbf{W})$ as $\lambda^q \rightarrow \infty$.

Intuitively, when $\lambda^q \rightarrow \infty$, we can roughly draw the second conclusion based on Eqn.

6.3. When $\lambda^q \rightarrow 0$, we have ²

$$\begin{aligned}
 NRAE_{p,q}(f(\mathbf{x}_i, \mathbf{W}), y_i) &= \frac{1}{\lambda^q} \log \frac{1}{m} \sum_{i=1}^m e^{\lambda^q \|f(\mathbf{x}_i, \mathbf{W}) - y_i\|^p} \\
 &= \frac{1}{\lambda^q} \log \left(1 + \frac{1}{m} \sum_{i=1}^m \lambda^q \|f(\mathbf{x}_i, \mathbf{W}) - y_i\|^p + O(\lambda^{2q}) \right) \\
 (5.19) \quad &= \frac{1}{m} \sum_{i=1}^m \|f(\mathbf{x}_i, \mathbf{W}) - y_i\|^p
 \end{aligned}$$

Please refer to the supplemental materials for the proofs. More rigid proofs that can be generalized to L_p -norm error are also given in (Lo 2010). In SimNets, the authors also include quite similar discussions about the robustness with respect to L_p -norm error (Cohen & Shashua 2014).

5.4 Learning Methods

We propose a novel learning method to training DNNs with NRAE, called the Adaptive Normalized Risk-Avering Training (ANRAT) approach. Instead of manually tuning λ like GDC (Lo, Gui, & Peng 2012), we learn λ adaptively in error backpropagation by considering λ as a parameter instead of a hyperparameter. The learning procedure is standard batch SGD. We show it works quite well in theory and practice.

The loss function of ANRAT is

$$(5.20) \quad l(W, \lambda) = \frac{1}{\lambda^q} \log \frac{1}{m} \sum_{i=1}^m e^{\lambda^q \|f(\mathbf{x}_i, \mathbf{W}) - y_i\|^p} + a \|\lambda\|^{-r}$$

²Consider the rules $e^x = 1 + x + \frac{x^2}{2} + \dots (x \rightarrow 0)$ and $\log x = x - \frac{x^2}{2} + \dots (x \rightarrow 0)$

Except for NRAE, we use a penalty term $a||\lambda||^{-r}$ to control the changing rate of λ . While minimize the NRAE score, small λ is penalized to regulate the convexity region. a is a hyperparameter to control the penalty index. The first-order derivatives on weight and λ are

$$(5.21) \quad \frac{dl(W, \lambda)}{dW} = \frac{p \sum_{i=1}^m e^{\lambda^q \alpha_i(W)^{p-1}} \frac{\partial f(\mathbf{x}_i, \mathbf{W})}{\partial \mathbf{W}}}{\sum_{i=1}^m e^{\lambda^q \alpha_i(W)^{p-1}}}$$

$$(5.22) \quad \frac{dl(W, \lambda)}{d\lambda} = \frac{-q}{\lambda^{q+1}} \log \frac{1}{m} \sum_{i=1}^m e^{\lambda^q \alpha_i(W)^p}$$

$$(5.23) \quad + \frac{q}{\lambda} \frac{\sum_{i=1}^m e^{\lambda^q \alpha_i(W)^p} \alpha_i(W)^p}{\sum_{i=1}^m e^{\lambda^q \alpha_i(W)^p}}$$

$$(5.24) \quad - ar \lambda^{-r-1}$$

We make a transformation on Eqn. 5.24 to better understand the gradient with respect to λ . Note that $k_i = \frac{e^{\lambda^q \alpha_i(W)^p}}{\sum_{i=1}^m e^{\lambda^q \alpha_i(W)^p}}$ is actually performing like a probability ($\sum_{i=1}^m k_i = 1$). Ignoring the penalty term, Eqn. 5.24 can be formulated as follows:

$$(5.25) \quad \begin{aligned} \frac{dl(W, \lambda)}{d\lambda} &= \frac{q}{\lambda} \left(\sum_{i=1}^m k_i \alpha_i(\mathbf{W})^p - NRAE \right) \\ &= \frac{q}{\lambda} (\mathbb{E}(\alpha(\mathbf{W})^p) - NRAE) \\ &\approx \frac{q}{\lambda} (L_P\text{-norm error} - NRAE) \end{aligned}$$

Note that as $\alpha_i(\mathbf{W})^p$ becomes smaller, the expectation on $\alpha_i(\mathbf{W})^p$ approaches the standard L_p -norm error. Thus, the gradient on λ is approximately the difference between NRAE and the standard L_p -norm error. Because large λ can incur plateaus to prevent NRAE from finding better optima using batch SGD (Lo, Gui, & Peng 2012), they need GDC to gradually deconvexify the NRAE to make the error space well shaped and stable. Through Eqn. 6.11, ANRAT solve this problem in a more flexible and adaptive manner.

When NRAE is larger, Eqn. 6.11 remains negative and makes λ increase to enlarge the convexity region, facilitating the search in the error space for better optima. When NRAE is smaller, the learned parameters are seemingly going through the optimal "tunnel" for better optima. Eqn. 6.11 becomes positive to decrease λ and helps NRAE not deviate far from the manifold of the standard L_p -norm error to make the error space stable without large plateaus. Thus, ANRAT adaptively adjusts the convexity index to find an optimal trade-off between better solutions and stability.

This training approach has more flexibility. The gradient on λ as the weighted difference between NRAE and the standard L_p -norm error, enables NRAE to approach the L_p -norm error by adjusting λ gradually. Intuitively, it keeps searching the error space near the manifold of the L_p -norm error to find better optima in a way of competing with and at the same time relying on the standard L_p -norm error space.

In Eqn. 5.20, the penalty weight a and index r control the convergence speed by penalizing small λ . Smaller a emphasizes tuning λ to allow faster convergence speed between NRAE and L_p -norm error. Larger a forces larger λ for a better chance to find a better local optimum but runs the risk of plateaus and deviating far from the stable error space. r regulates the magnitude of λ and its derivatives in gradient descent.

5.5 Experiments

We present the results from a series of experiments designed on the MNIST and CIFAR-10 datasets to test the effectiveness of ANRAT for visual recognition with DNNs. We did not explore the full hyperparameter in Eqn. 5.20. Instead we fix the hyperparameter at $p = 2$, $q = 2$ and $r = 1$ to mainly compare with MSE. So the final loss function of ANRAT we optimized is

$$(5.26) \quad l(W, \lambda) = \frac{1}{\lambda^2} \log \frac{1}{m} \sum_{i=1}^m e^{\lambda^2 \|f(\mathbf{x}_i, \mathbf{W}) - y_i\|^2} + a|\lambda|^{-1}$$

This loss function is minimized by batch SGD without complex methods, such as momentum, adaptive/hand tuned learning rates or tangent prop. The learning rate and penalty weight a are selected in $\{1, 0.5, 0.1\}$ and $\{1, 0.1, 0.001\}$ on validation sets respectively. The initial λ is fixed at 10. We use the hold-out validation set to select the best model, which is used to make predictions on the test set. All experiments are implemented quite easily in Python and Theano to obtain GPU acceleration (Bastien *et al.* 2012).

The MNIST dataset (LeCun *et al.* 1998) consists of hand written digits 0-9 which are 28x28 in size. There are 60,000 training images and 10,000 testing images in total. We use 10000 images in training set for validation to select the hyperparameters and report the performance on the test set. We test our method on this dataset without data augmentation.

The CIFAR-10 dataset (Krizhevsky & Hinton 2009) is composed of 10 classes of natural images. There are 50,000 training images in total and 10,000 testing images. Each image is an RGB image of size 32x32. For this dataset, we adapt pylearn2 (Goodfellow *et al.* 2013a) to apply the same global contrast normalization and ZCA whitening as was used by Goodfellow *et. al* (Goodfellow *et al.* 2013b). We use the last 10,000 images of the training set as validation data for hyperparameter selection and report the test accuracy.

5.6 Results and Discussion

5.6.1 Results on ConvNets

On the MNIST dataset we use the same structure of LeNet5 with two convolutional max-pooling layers but followed by only one fully connected layer and a densely connected

softmax layer. The first convolutional layer has 20 feature maps of size 5×5 and max-pooled by 2×2 non-overlapping windows. The second convolutional layer has 50 feature maps with the same convolutional and max-pooling size. The fully connected layer has 500 hidden units. An l_2 prior was used with the strength 0.05 in the Softmax layer. Trained by ANRAT, we can obtain a test set error of 0.52%, which is the best result we are aware of that does not use dropout on the pure ConvNets. We summarize the best published results on the standard MNIST dataset in Table 6.2.

The best performing neural networks for pure ConvNets that does not use dropout or unsupervised pretraining achieve an error of about 0.69% (Ngiam *et al.* 2011). They demonstrated this performance with L-BFGS. Using dropout, ReLU and a response normalization layer, the error reduces to 0.55% (Zeiler & Fergus 2013). Prior to that, Jarrett *et al.* showed by increasing the size of the network and using unsupervised pretraining, they can obtain a better result at 0.53% (Jarrett *et al.* 2009). Previous state of the art is 0.47% (Goodfellow *et al.* 2013b) for a single model without on the original MNIST dataset. Using batch SGD to optimize either CE or MSE on the ConvNets described above, we can get an error rate at 0.93%. Replacing the training methods with ANRAT using batch SGD leads to a sharply decreased validation error of 0.66% with a test error at 0.52%. With dropout and ReLU the error rate drops to 0.39%, which is the same with the best results without averaging or data augmentation (Table 6.2) but we only use standard Convnets and simple experimental settings.

Fig. 5.1 (a) shows the progression of training, validation and test errors over 160 training epochs. The errors trained on MSE plateau as it can not train the ConvNets sufficiently and seems like underfit. Using ANRAT, the validation and test errors remain decreasing

³(1)(Mairal *et al.* 2014);(2)(Lee *et al.* 2014);(3)(LeCun *et al.* 1998);(4)(Ranzato *et al.* 2007);(5)(Ngiam *et al.* 2011);(6)(Ranzato *et al.* 2007);(7)(Poultney *et al.* 2006);(8)(Zeiler & Fergus 2013);(9)(Jarrett *et al.* 2009)

Method ³	Error %
Convolutional Kernel Networks + L-BFGS-B ⁽¹⁾	0.39
Deeply Supervised Nets + dropout ⁽²⁾	0.39
ConvNets (Lenet-5) ⁽³⁾	0.95
ConvNets + MSE/CE (this section)	0.93
large ConvNets, random feature ⁽⁴⁾	0.89
ConvNets + L-BFGS ⁽⁵⁾	0.69
large ConvNets, unsup pretraining ⁽⁶⁾	0.62
ConvNets, unsup pretraining ⁽⁷⁾	0.6
ConvNets + dropout ⁽⁸⁾	0.55
large ConvNets, unsup pretraining ⁽⁹⁾	0.53
ConvNets + ANRAT (This paper)	0.52
ConvNets + ANRAT + dropout (This paper)	0.39

Table 5.1. Test set misclassification rates of the best methods that utilized convolutional networks on the original MNIST dataset using single model.

along with the training error. During training, λ sharply decrease, regulating the tunnel of NRAE to approach the manifold of MSE. Afterward the penalty term becomes significant, force λ to grow gradually while expanding the convex region for higher chance to find the better optimum (Figure 5.1 (b)).

Using the same network architecture described above, we trained two ConvNets using MSE and ANRAT respectively and compare their performance. Fig. 5.1 (a) shows the progression of train, validation and test errors over 160 training epochs. The errors trained on MSE plateau as it can not train the ConvNets sufficiently and underfits. Using ANRAT, the validation and test errors remain decreasing along with the training error. During training, λ sharply decrease, regulating the tunnel of NRAE to approach the manifold of MSE. Afterward the penalty term becomes significant, force λ to grow gradually while expanding the convex region for higher chance to find the better optimum (Fig. 5.1 (b)).

Our next experiment is performed on the CIFAR-10 dataset. We observed significant overfitting using both MSE and ANRAT with the fixed learning rate and batch SGD, so

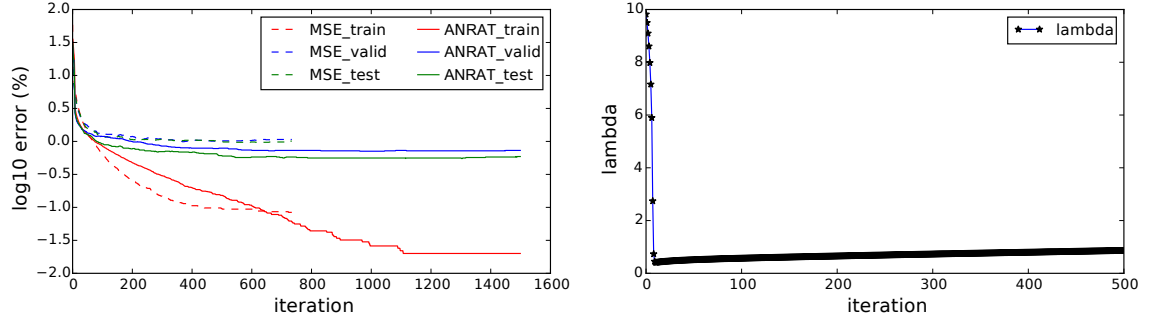


FIG. 5.1. (a). MNIST train, validation and test error rates throughout training with Batch SGD for MSE and ANRAT with l_2 priors. We cross-validated the learning rate and regularization weight on validation set for ConvNets (left). (b). The curve of λ throughout ANRAT training (right).

dropout is applied to prevent the co-adaptation of weights and improve generalization. We use a similar network layout as in (Srivastava *et al.* 2014) but with only two convolutional max-pooling layers. The first convolutional layer has 96 feature maps of size 5×5 and max-pooled by 2×2 non-overlapping windows. The second convolutional layer has 128 feature maps with the same convolutional and max-pooling size. The fully connected layer has 500 hidden units. Dropout was applied to all the layers of the network with the probability of retaining a hidden unit being $p = (0.9, 0.75, 0.5, 0.5, 0.5)$ for the different layers of the network. Using batch SGD to optimize CE on the simple configuration of ConvNets + dropout, a test accuracy of 80.6 % is achieved (Krizhevsky, Sutskever, & Hinton 2012). We also reported the performance at 80.58% with MSE instead of CE with the similar network layout. Replacing the training methods with ANRAT using batch SGD gives a test accuracy of 85.15%. This is superior to the results obtained by MSE/CE and unsupervised pretraining. In Table. 5.2, our result with simple setting is shown to be competitive to those achieved by different ConvNet variants.

⁴(1)(Zeiler & Fergus 2013);(2)(Srivastava *et al.* 2014);(3)(Goodfellow *et al.* 2013b);(4)(Zeiler & Fergus 2013);(5)(Coates & Ng 2011);(6)(Min Lin 2014);(7)(Lee *et al.* 2014)

Table 5.2. Test accuracy of the best methods that utilized convolutional framework on CIFAR-10 dataset without data augmentation.

Method ⁴	Acc %
ConvNets + Stochastic pooling + dropout ⁽¹⁾	84.87
ConvNets + dropout + Bayesian hyperopt ⁽²⁾	87.39
ConvNets + Maxout + dropout ⁽³⁾	88.32
Convolutional NIN + dropout ⁽⁶⁾	89.6
Deeply Supervised Nets + dropout ⁽⁷⁾	90.31
ConvNets + MSE + dropout (this section)	80.58
ConvNets + CE + dropout ⁽⁴⁾	80.6
ConvNets + VQ unsup pretraining ⁽⁵⁾	82
ConvNets + ANRAT + dropout (This section)	85.15

5.6.2 Results on Multilayer Perceptron

While ConvNets leads the state of the art performance on visual recognition, general Multilayer Perceptron (MLP) is an important framework for generative and discriminative learning. Unsupervised layer-wise pretraining by RBM or auto-encoder demonstrate superb performance boosting on visual recognition on MLP.

On the MNIST dataset, MLPs with unsupervised pretraining has been well studied in recent years, so we select this dataset to compare ANRAT in shallow and deep MLPs with MSE/CE and unsupervised pretraining. For the shallow MLPs, we follow the network layout as in (Gui, Lo, & Peng 2014; LeCun *et al.* 1998) that has only one hidden layer with 300 neurons. We build the stacked architecture and deep network using the same architecture as (Larochelle *et al.* 2009) with 500, 500 and 2000 hidden units in the first, second and third layers, respectively. The training approach is purely batch SGD with no momentum or adaptive learning rate. No weight decay or other regularization technique is applied in our experiments.

Experiment results in Table. 5.3 show that the deep MLP classifier trained by the

ANRAT method has the lowest test error rate (1.45%) of benchmark MLP classifiers with MSE/CE under the same settings. It indicates that ANRAT has the ability to provide reasonable solutions with different initial weight vectors. This result is also better than deep MLP + supervised pretraining or Stacked Logistic Regression networks. We note that the deep MLP using unsupervised pretraining (auto-encoders or RBMs) remains to be the best with test error at 1.41% and 1.2%. Unsupervised pretraining is effective in initializing the weights to obtain a better local optimum. Compared with unsupervised pretraining + fine tuning, ANRAT sometimes still fall into the slightly worse local optima in this case. However, ANRAT is significantly better than MSE/CE without unsupervised pretraining.

Experimental results in Table. 5.3 show that the deep MLP classifier trained by ANRAT method has the lowest test error rate (1.55%) than benchmark MLP classifiers with MSE/CE with the same settings. It indicates that the ANRAT method has the ability to provide reasonable generalization results with different initial weight vectors. This result are also better than deep MLP + supervised pretraining or Stacked Logistic Regression network. However, we note that the deep MLP using unsupervised pretraining (auto-encoders or RBM) is remaining the best with the test error at 1.41% and 1.2%. Unsupervised pretraining is still effective to better initialize the weight to obtain a better local optimum. Compared with unsupervised pretraining + fine tuning, ANRAT sometimes still fall into the slightly worse local optimum in this case. However, ANRAT is significantly better than MSE/CE without unsupervised pretraining.

Interestingly, we do not observe significant advantages with ANRAT in shallow MLPs. Although in early literature, the error rate on shallow MLPs were reported as 4.7% (LeCun *et al.* 1998) and 2.7% with GDC (Gui, Lo, & Peng 2014), both recent papers using CE (Larochelle *et al.* 2009) and our own experiments with MSE can achieve error rate of 1.93% and 2.02%, respectively. Trained by ANRAT, we can have a test rate at 1.94%. This performance is slightly better than MSE, but it is statistically identical to the performance

obtained by CE.⁵ One possible reason is that in shallow networks which can be trained quite well by standard back propagation with normalized initializations, the local optimum achieved with MSE/CE is quite nearly a global optimum or good saddle point. Our result is also corresponding to the conclusion in (Dauphin *et al.* 2014), in which Dauphin et al. extend previous findings on networks with a single hidden layer to show theoretically and empirically that most badly suboptimal critical points are saddle points. Even with better convexity property, ANRAT is as good as MSE/CE in shallow MLPs. However, we find that the problem of poor local optimum becomes more manifest in deep networks. It is easier for ANRAT to find a way towards the better optimum near the manifold of MSE. For the sake of space, please refer to supplemental materials for the results on the shallow Denoised Auto-encoder. The conclusion is consistent that ANRAT performs better when attacking more difficult learning/fitting problems. While ANRAT is slightly better than CE/MSE + SGD on DA with uniform masking noise, it achieves a significant performance boost when Gaussian block masking noise is applied. Moreover, the exponential form of NRAE helps to alleviate vanishing of gradient also help training deep networks.

5.6.3 Results on Denoised Auto-encoder

Our last experiments focus on generative learning instead of discriminative learning. In (Vincent *et al.* 2010), Denoised Auto-encoder (DA) is well studied to show the ability to learn useful representations. We were considering if ANRAT is likely to lead to the learning of feature detectors that detect better structure in the input patterns with more sufficient training.

On MNIST dataset, we use a single-layer DA with the local masking noise at the

⁵in (Larochelle *et al.* 2009), the author do not report their network settings of the shallow MLP + CE, which may differ from 784-300-10.

⁶(1)(Larochelle *et al.* 2009);(2)(LeCun *et al.* 1998);(3)(Gui, Lo, & Peng 2014)

Table 5.3. Test error rate of deep/shallow MLP with different training techniques.

Method ⁶	Error %
Deep MLP + supervised pretraining ⁽¹⁾	2.04
Stacked Logistic Regression Network ⁽¹⁾	1.85
Stacked Auto-encoder Network ⁽¹⁾	1.41
Stacked RBM Network ⁽¹⁾	1.2
Shallow MLP + MSE ⁽²⁾	4.7
Shallow MLP + GDC ⁽³⁾	2.7 ± 0.03
Shallow MLP + MSE (this section)	2.02
Shallow MLP + ANRAT (this section)	1.94
Shallow MLP + CE ⁽¹⁾	1.93
Deep MLP + CE ⁽¹⁾	2.4
Deep MLP + MSE (this section)	1.91
Deep MLP + ANRAT (this section)	1.45

levels from 30% to 60%. That is, we randomly set the pixels to be 0 evenly. Another noise paradigm is Gaussian block noise. We randomly select centroids to span masking blocks in terms of each centroid with a Gaussian distribution. The noise level differs from 33.3% to 52.6% accordingly. Figure. 5.2 shows three examples of different corruption level. We train the DA using ANRAT and MSE with pure batch SGD with fixed learning rate among $\{1, 0.1, 0.01\}$ for 1000 epochs and report the best reconstruction MSE.

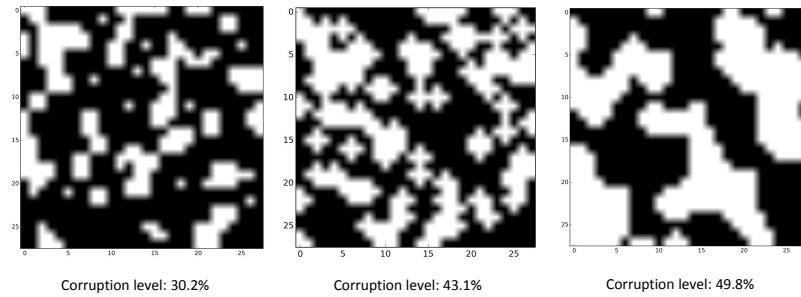


FIG. 5.2. Example maps of Gaussian block masking.

As in the shallow network, we found the advantage of ANRAT is statistically significant in DA but the difference is very small when using masking noise (Figure. 5.4). Although masking makes it harder to reconstruct original images, the strong spatial correlations are still existed between masked pixels and the "good" pixels nearby, thus helps DA to learn the generative distribution from the conditional probability with respect to the manifold of the added noise. The local optimum found by both SGD training with MSE and ANRAT is nearly optimal. When applying Gaussian block noise as shown in Figure. 5.2, masking blocks cover several patches instead of scattered points. It is much harder to learn the spatial correlation within the masking areas. On this tougher task, ANRAT performs much better than using batch SGD on MSE. In Figure. 5.4, the reconstruction MSE using ANRAT on each image is 1.078, which is 24.09% lower than 1.432 obtained by batch SGD on MSE.

Table 5.4. Reconstruction MSE of DA output trained by ANRAT and MSE error criterion at different masking level using random masking or Gaussian block masking. The mean and standard deviation over 20 times running are reported.

Corruption Level	MSE	ANRAT
30%	4.552 (0.006)	4.348 (0.005)
40%	5.697 (0.007)	5.57 (0.005)
50%	7.112 (0.005)	7.071 (0.007)
60%	9.245 (0.008)	9.1 (0.008)
Gaussian Isotropic Masking (33.3% - 52.6%)	1.432 (0.0002)	1.078 (0.009)

In Figure. 5.3 (a), the curve of training error using ANRAT penetrates across the curve trained by MSE after several epochs and is consistently staying lower. Note that NRAE and L_P norm can be compared since they are homeomorphism and in the same magnitude, but there is still a "gap" between them. So, λ is still forced to fluctuate, gradually grow larger to find more optimal tunnel towards the better optimum (Figure 5.3 (b)). This keeps ANRAT searching for better results near the manifold of MSE through a gradual convexing

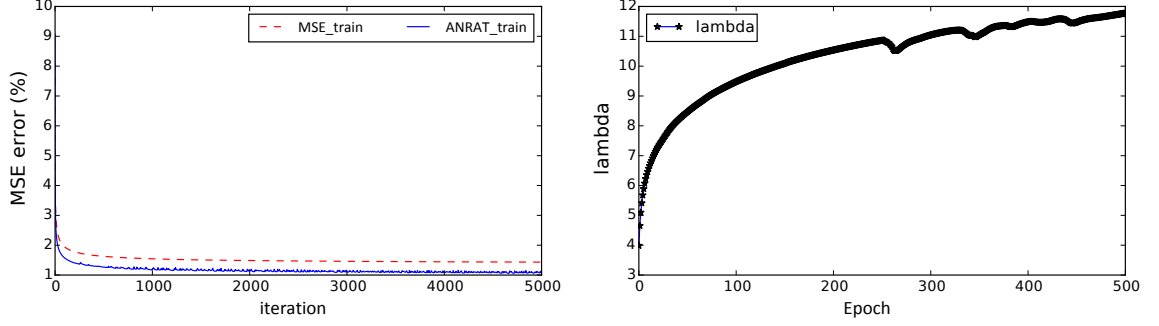


FIG. 5.3. (a). Reconstruction MSE throughout training by Batch SGD for MSE and ANRAT (left). Gaussian block noise is applied at input layer. (b). The curve of λ throughout ANRAT (right).

on its convexity region. In Figure. 5.4, DA using ANRAT learned a lightly more diverse features than SGD on MSE when uniform masking is applied. In the case of Gaussian block masking, DA learned more clear meaningful features with the help of ANRAT.

5.7 Conclusions and Outlook

In this section, we introduce a novel approach, Adaptive Normalized Risk-Averting Training (ANRAT), to help train deep neural networks. Theoretically, we prove the effectiveness of Normalized Risk-Averting Error on its arithmetic bound, global convexity and local convexity lower-bounded by standard L_p -norm error when convexity index $\lambda \geq 1$. By analyzing the gradient on λ , we explained the reason why using back propagation on λ works. The experiments on deep/shallow network layouts demonstrate comparable or better performance with the same experimental settings among pure ConvNets and MLP + batch SGD on MSE and CE (with or without dropout). Other than unsupervised pre-training, it provides a new perspective to address the non-convex optimization strategy in DNNs. Theoretically, we prove the effectiveness of Normalized Risk-Averting Error on its arithmetic bound, global convexity and local convexity lower-bounded by standard L_p -

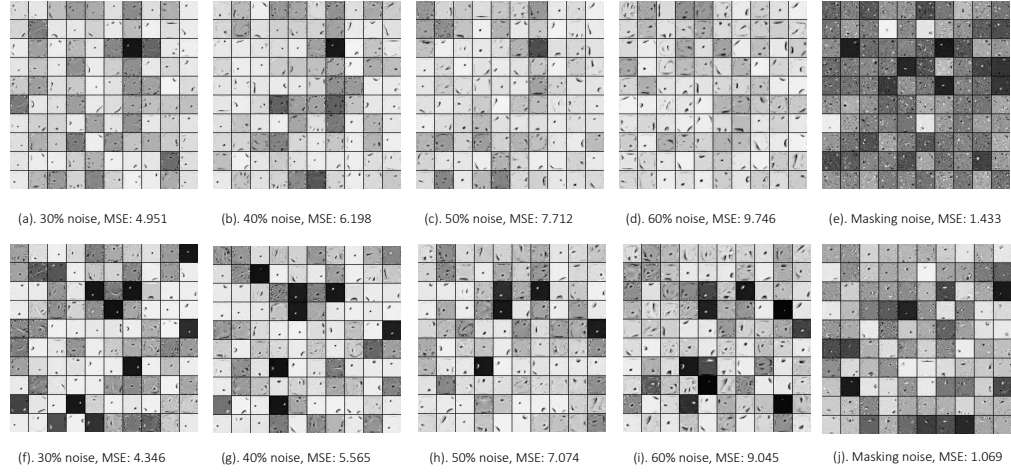


FIG. 5.4. 100 feature maps learned by Denoised Auto-encoder using ANRAT and batch SGD on MSE. Different corruption level of uniform masking and Gaussian block masking are applied at the input layer.

norm error when convexity index $\lambda \geq 1$. By analyzing the gradient on λ , we explained the reason why using back propagation on λ works.

Empirically, we compare ANRAT with batch SGD training on MSE/CE and unsupervised pretraining using pure ConvNets on MNIST and CIFAR-10 datasets. ANRAT achieves comparable or better performance with the same experiment settings among those use pure ConvNets + batch SGD and/or dropout. An overview of the comparisons among the best results obtained from different training methods and model variants is also provided to validate the performance of ANRAT is among the state of the art. Except for ConvNets, we evaluate ANRAT in shallow/deep multilayer perceptrons. Although the advantage of ANRAT is not significant against CE/MSE + SGD in shallow neural networks, it overtakes CE/MSE + SGD in deep neural networks and approaches to the performance achieved by unsupervised pretraining such as RBM and auto-encoder. To better understand ANRAT on generative models, experiments on shallow Denoised Auto-encoders are also performed

with uniform masking noise and Gaussian block masking noise at the input layer. While ANRAT is slightly better than CE/MSE + SGD on DA with the uniform masking noise, it achieves significant performance boosting when Gaussian block masking noise is applied.

Finally, while these early results are very encouraging, clearly further research is warranted to address the questions that arise from non-convex optimization in deep neural networks. It is preliminarily showed that in order to generalize to a wide array of tasks, unsupervised and semi-supervised learning using unlabeled data is crucial. One interesting future work is to take advantage of unsupervised/semi-supervised pretraining with the non-convex optimization methods to train deep neural networks by finding the nearly global optimum. Another crucial question is to guarantee the generalization capability by preventing overfitting. Finally, we are quite interested in generalizing our approach to recurrent neural networks. We leave as future work any performance improvement on benchmark datasets by considering the cutting-edge approach to improve training and generalization performance such as Bayesian hyperparameter optimization, network layout variants, Maxout or SFN, etc.

Chapter 6

ROBUST STATISTICS — ADAPTIVE NORMALIZED ANOMALY-AVERTING TRAINING FOR NEURAL NETWORKS

6.1 Introduction

So far, we only focused on the situations when $\lambda > 0$. However, when $\lambda < 0$ and is small enough, from Equation. 6.7 the Hessian matrix is still definite positive to guarantee the error function is convex. The credits should be given to Lo who firstly proposed such an idea in (Lo & Bassu 2001) and we consistently investigate Lo’s early work with different naming systems. When $\lambda \rightarrow -\infty$, we called the error estimator Anomaly-Averting Estimator (AAE, in Lo’s work, he called it Risk-seeking Error). AAE not only ensures the convexity but also retains high robustness to outliers. Note that when λ is largely negative, we still have the register underflow problem. The same normalized pipeline in the previous section also works on AAE. In this section, we will explore the optimality and robustness of Normalized Anomaly-Averting Estimator (NAAE) by pushing λ to $-\infty$. A more specific theoretical statements are provided in l_2 norm to clarify its connection with square loss. Robustness is widely required in control theory and dynamical systems design. Thus, we also evaluate NAAE on both function approximation problems and machine learning tasks.

6.2 Background

Function approximation has many applications in science and engineering, such as machine learning, pattern recognition, signal processing and control theory. As universal approximators (Hornik, Stinchcombe, & White 1989), neural networks (NNs) often deliver very good performance and have become the standard for several machine learning tasks given their recent success in various applications (Jaderberg *et al.* 2015; Lee *et al.* 2014; Szegedy *et al.* 2015; Karpathy & Fei-Fei 2015).

Error-free data are rarely provided in applications. Instead, data are usually contaminated by noise and outliers. Noise reflects inaccuracies in observations and the stochastic nature of the underlying process. NNs deal with such noise quite efficiently by optimizing the minimum squared error (MSE) or cross entropy (CE) between the observed and predicted values/labels with ℓ_1/ℓ_2 regularization, sparsity or adversary resistant learning (Sra, Nowozin, & Wright 2012). Recent developments in computer vision take advantage of manually added noise and distortion to further enhance the data (Krizhevsky, Sutskever, & Hinton 2012; Gan *et al.* 2015) or to capture reverse conditional probability within generative models (Bengio *et al.* 2013; Goodfellow *et al.* 2014). Outliers can be arbitrary, unbounded, and not from any specific distribution, reflecting sudden abnormal changes or a mixture of rare but different phenomena. It has been noted that outliers in routine data are infrequent but heavily impact even high dimensional data (Aggarwal & Yu 2001; Kriegel, Zimek, & others 2008). When fitting a model or learning an objective function with rare but significant outliers whose distribution and impact are both unknown, they need to be identified and eliminated to get rid of their effect. Otherwise, the model overfits or underfits easily if training procedure is sufficient or not. For example, in the image/video classification task, some images or videos may be corrupted unexpectedly due to the error of sensors or severe occlusions of objects. Such outliers can skew parameter estimation

severely and hence destroy the performance of the learned model. Alternatively, outliers are sometimes supposed to be examined closely, as they may be of interest themselves given the sample applications of intrusion or fraud detection.

The maximal likelihood principle is not robust against outliers among different approximation and learning models like logistic regression (LR) and (Feng *et al.* 2014) generalized linear models (Künsch, Stefanski, & Carroll 1989). High breakdown methods have been developed including least median of squares (Rousseeuw 1984), least trimmed squares (Rousseeuw 1984) and most recently trimmed inner product methods for linear regression. Several works have investigated multiple approaches to robustify LR (Pregibon 1982; Stefanski, Carroll, & Ruppert 1986; Tibshirani & Manning 2013). The majority of them are M-estimator based: minimizing a complicated and more robust loss function than the standard loss function (negative log-likelihood). A robust backpropagation algorithm is proposed to optimize a robust estimator derived from the M-estimator to train NNs (Chen & Jain 1994). (Liano 1996) used M-estimators to study the mechanism by which outliers affect the resulting NNs. The majority of the above work that is M-estimator based needs to predefine a threshold for determining the degree of contaminated data to achieve robustness to outliers. Besides, few of them discuss optimality while achieving the robustness. Here, by optimality, we refer to its ability to find a nearly global or better local optimum while achieving robustness to outliers by getting rid of large deviation among the data.

Our work is largely inspired by the exponentialized error criteria (Jacobson 1973; Whittle 1990) and its recent variations and modifies exponentialized error with two different training methods (Lo 2010; Gui, Lo, & Peng 2014; Wang, Oates, & Lo 2015). Our main contribution is to implement with our adaptive training approach and verify the idea that generalizing this estimator by pushing the robust-optimal (RO) index λ to $-\infty$ to obtain the robustness to the largely deviated outliers while preserving the optimality by the expansion of the convexity regions in the Hessian matrix. As a general error estimator, we

provide a quantitative analysis and validate its effectiveness on three function fitting and one visual recognition tasks.

6.3 Normalized Anomaly-Averting Estimator

Given standardized training samples $\{\mathbf{X}, y\} = \{(\mathbf{x}_1, y_1), (\mathbf{x}_2, y_2), \dots, (\mathbf{x}_m, y_m)\}$, $f(\mathbf{x}_i, \mathbf{W})$ is the learned model with parameters $\mathbf{W}$. The loss function of mean squared loss (MSE) is defined as:

$$l_2(f(\mathbf{x}_i, \mathbf{W}), y_i) = \frac{1}{m} \sum_{i=1}^m (f(\mathbf{x}_i, \mathbf{W}) - y_i)^2$$

As a modified exponentialized error of MSE, the Anomaly-Averting Estimator (AAE) is defined as

$$(6.1) \quad AAE = \frac{1}{m} \sum_{i=1}^m e^{\lambda(f(\mathbf{x}_i, \mathbf{W}) - y_i)^2} \quad (\lambda < 0)$$

The only difference between the AA estimator and RAE in (Lo 2010) is that λ stays negative, which makes it serve as the robust-optimal (RO) index to control both the degree of optimality and robustness to outliers. Note that when $\lambda \rightarrow -\infty$, it is not bounded and suffers from less stability, exponential magnitude and arithmetic overflow when using gradient descent in implementations. We define the Normalized Anomaly-Averting Estimator (NAAE) as:

$$(6.2) \quad \begin{aligned} NAAE(f(\mathbf{x}_i, \mathbf{W}), y_i) &= \frac{1}{\lambda} \log AAE(f(\mathbf{x}_i, \mathbf{W}), y_i) \\ &= \frac{1}{\lambda} \log \frac{1}{m} \sum_{i=1}^m e^{\lambda(f(\mathbf{x}_i, \mathbf{W}) - y_i)^2} \quad (\lambda < 0) \end{aligned}$$

6.3.1 Robustness

NAAE is a bounded estimator. The bounds exert not only the stable property as an estimator, but also an interesting dynamics to control the degree of robustness.

Theorem 5. *NAAE($f(\mathbf{x}_i, \mathbf{W}), y_i$) is bounded between the minimal and maximal squared errors.*

Proof. Let

$$\begin{aligned}\alpha_i &= f(\mathbf{x}_i, \mathbf{W}) - y_i \\ \alpha_{min}^2 &= \min\{\alpha_i^2, i = 1, \dots, m\} \\ \alpha_{max}^2 &= \max\{\alpha_i^2, i = 1, \dots, m\} \\ \beta_i &= e^{\lambda(\alpha_i^2 - \alpha_{min}^2)}\end{aligned}$$

Then we can re-write Eqn. 6.2 as

$$\begin{aligned}NAAE &= \frac{1}{\lambda} \log \frac{1}{m} e^{\lambda \alpha_{min}^2} \sum_{i=1}^m \beta_i \\ (6.3) \quad &= -\frac{1}{\lambda} \log m + \alpha_{min}^2 + \frac{1}{\lambda} \log \sum_{i=1}^m \beta_i\end{aligned}$$

let $G = \alpha_{max}^2 - \alpha_{min}^2$ to be the maximal margin of the sample squared errors, considering $\lambda < 0$, $e^{\lambda G} \leq \beta_i \leq 1$, so

$$\frac{1}{\lambda} \log m \leq \frac{1}{\lambda} \log \sum_{i=1}^m \beta_i \leq \frac{1}{\lambda} \log m + G$$

We can derive the bounds of NAAE as

$$\alpha_{min}^2 \leq NAAE \leq \alpha_{max}^2$$

□

The above conclusions indicate that squared errors regulate the bounds of NAAE. The next theorem shows that λ controls NAAE to perform interactively with MSE and its relationship to the minimin operator.

Theorem 6. *when $\lambda \rightarrow -\infty$, $NAAE(f(\mathbf{x}_i, \mathbf{W}), y_i)$ approaches to an minimin operator; if $\lambda \rightarrow 0$, $NAAE(f(\mathbf{x}_i, \mathbf{W}), y_i) \rightarrow MSE$.*

Proof. Given $\lambda < 0$, we consider the objective function

$$W = \arg \min_{\mathbf{W}} NAAE(f(\mathbf{x}_i, \mathbf{W}), y_i)$$

Considering equation 6.3, it is easy to see that when $\lambda \rightarrow -\infty$, the objective function becomes

$$W = \arg \min_{\mathbf{W}} \alpha_{min}^2$$

So a large λ controls the NAAE to approach a minimin operator. Next we consider the

situation when $\lambda \rightarrow 0$ ¹.

$$\begin{aligned}
 & NAAE(f(\mathbf{x}_i, \mathbf{W}), y_i) \\
 &= \frac{1}{\lambda} \log \frac{1}{m} \sum_{i=1}^m e^{\lambda(f(\mathbf{x}_i, \mathbf{W}) - y_i)^2} \\
 &= \frac{1}{\lambda} \log \left(1 + \frac{1}{m} \sum_{i=1}^m \lambda(f(\mathbf{x}_i, \mathbf{W}) - y_i)^2 + O(\lambda^2) \right) \\
 (6.4) \quad &= \frac{1}{m} \sum_{i=1}^m (f(\mathbf{x}_i, \mathbf{W}) - y_i)^2
 \end{aligned}$$

□

More rigid proofs that can be generalized to L_p -norm error are given in (Lo 2010).

Equations 6.3 and 6.4 explain how λ controls the robustness level. When $\lambda \rightarrow 0$, NAAE approaches MSE with a breakdown point of 0%, meaning that a single observation can change it arbitrarily. Further, it is highly influenced by outliers. When $\lambda \rightarrow -\infty$, NAAE performs like a minimin operator to address only the minimal error. Outliers can be eliminated if they largely deviate from the objective space, or say manifold, so that NAAE concentrates on the smaller errors and ignores the impact of those large errors. But it is still vulnerable to noise and outliers that are much closer to the objective manifold than the majority of the clean data.

When λ changes between 0 and $-\infty$, the exponential sum of the squared error offsets the minimin operator to further address the situations when the outliers are close to the objective manifold. By adjusting λ , NAAE is able to handle both large deviations and false positive samples, which helps to achieve robustness to different type of outliers.

¹Consider the rules $e^x = 1 + x + \frac{x^2}{2} + \dots (x \rightarrow 0)$ and $\log x = x - \frac{x^2}{2} + \dots (x \rightarrow 0)$

6.3.2 Optimality

To prove the optimality of NAAE, we start with AAE and generalize our conclusion to NAAE. For clarity, we write the matrix derivatives in functional form. In our proof we use the quadratic form. The Jacobian matrix of Eqn. 6.1 is

$$(6.5) \quad \begin{aligned} J(\mathbf{W}) = & \frac{2\lambda}{m} \sum_{i=1}^m e^{\lambda(f(\mathbf{x}_i, \mathbf{W}) - y_i)^2} \\ & \times (f(\mathbf{x}_i, \mathbf{W}) - y_i) \frac{\partial f(\mathbf{x}_i, \mathbf{W})}{\partial \mathbf{W}} \quad (\lambda < 0) \end{aligned}$$

The Hessian matrix of Eqn. 6.1 is

$$(6.6) \quad H(\mathbf{W}) = \frac{2}{m} \sum_{i=1}^m e^{\lambda(f(\mathbf{x}_i, \mathbf{W}) - y_i)^2} \left\{ \lambda \frac{\partial f(\mathbf{x}_i, \mathbf{W})}{\partial \mathbf{W}}^2 \right.$$

$$(6.7) \quad + 2\lambda^2 (y_i - f(\mathbf{x}_i, \mathbf{W}))^2 \frac{\partial f(\mathbf{x}_i, \mathbf{W})}{\partial \mathbf{W}}^2$$

$$(6.8) \quad + \lambda (y_i - f(\mathbf{x}_i, \mathbf{W})) \frac{\partial^2 f(\mathbf{x}_i, \mathbf{W})}{\partial \mathbf{W}^2} \}$$

We assume the training sample is standardized and the error space is a unit sphere $|\mathcal{R}^n \leq 1|$ (that is, $|y_i - f(\mathbf{x}_i, \mathbf{W})| < 1$). Because AAE has the sum-exponential form, its Hessian matrix is tuned exactly by the RO index λ . The following theorem indicates the relation between the convexity index and its convexity region.

Theorem 7. *Given the Anomaly-Averting Error criterion AAE, which is twice continuous differentiable. $J(\mathbf{W})$ and $H(\mathbf{W})$ are the corresponding Jacobian and Hessian matrix. As $\lambda \rightarrow \pm\infty$, the convexity region monotonically expands to the entire parameter space except for the subregion $S := \{W \in \mathcal{R}^n | \text{rank}(H(\mathbf{W})) < n, H(\mathbf{W}) < 0\}$.*

Proof. Both $\frac{\partial f(\mathbf{x}_i, \mathbf{W})}{\partial \mathbf{W}}^2$ and $\lambda^2 (y_i - f(\mathbf{x}_i, \mathbf{W}))^2$ are positive semi-definite, matrix 6.7 is semi-positive definite, but Eqn. 6.6 and 6.8 may be indefinite. Let $\alpha_i(\mathbf{W}) = f(\mathbf{x}_i, \mathbf{W}) - y_i$,

We rewrite Eqn. 6.7 in quadratic form:

$$\begin{aligned}
 &= p\lambda^2 \alpha_i(\mathbf{W})^T \frac{\partial f(\mathbf{x}_i, \mathbf{W})^T}{\partial \mathbf{W}} \cdot \frac{\partial f(\mathbf{x}_i, \mathbf{W})}{\partial \mathbf{W}} \alpha_i(\mathbf{W}) \\
 &= p\lambda^2 \alpha_i(\mathbf{W})^T Q \Lambda Q^T \alpha_i(\mathbf{W}) \\
 (6.9) \quad &= p\lambda^2 S(\mathbf{W})^T \Lambda S(\mathbf{W})
 \end{aligned}$$

$\Lambda = \text{diag}[\Lambda_1, \Lambda_2, \dots, \Lambda_m]$. 6.6 is positive semi-definite.

If $S(\mathbf{W})$ is a full-rank matrix, then Eqn. 6.9 is positive definite. When $\lambda \rightarrow \pm\infty$, the eigenvalues Λ becomes dominant in the leading principal minors (as well as eigenvalues) of the Hessian matrix to make $H(\mathbf{W})$ monotonically increasing with λ . Assume $P_\lambda := \{W \in \mathcal{R}^n | H(\mathbf{W}) > 0\}$, we have $P_{\lambda_1} \subset P_{\lambda_2}$ when $\lambda_1 < \lambda_2$. Considering Eqn. 6.6, 6.7 and 6.8 are all bounded, $\exists \psi(W)$, s.t. $H(\mathbf{W}) > 0$ when $|\lambda| > \psi(W)$.

When $S(\mathbf{W})$ is not a full-rank matrix, the determinant of all its $\binom{n}{k}$ submatrices is 0. Thus, in the subregion $S := \{W \in \mathcal{R}^n | \text{rank}(H(\mathbf{W})) < n, H(\mathbf{W}) < 0\}$, there is no parameters satisfied the convexity conditions ($\bigcup P_\lambda$). $\square$

Theorem 7 states that when the RO index λ decrease to infinity, the convexity region in the parameter space of AAE expands monotonically to the entire space except the intersection of a finite number of lower dimensional sets. The number of sets increases rapidly as the number m of training samples increases. Roughly speaking, large $|\lambda|$ and m cause the size of the convexity region to grow larger in the error space of AAE.

When $\lambda \rightarrow -\infty$, the error space can be perfectly stretched to be strictly convex, thus avoiding the local optimum to find a global optimum. The following theorem states the quasi-convexity of NAAE.

Theorem 8. *Given a parameter space $\{W \in \mathcal{R}^n\}$, Assume $\exists \psi(W)$, s.t. $H(\mathbf{W}) > 0$ when $|\lambda| > \psi(W)$ to guarantee the convexity of $AAE(f(\mathbf{x}_i, \mathbf{W}), y_i)$. Then, $NAAE(f(\mathbf{x}_i, \mathbf{W}), y_i)$*

is *quasi-convex* and share the same local and global optima with $AAE(f(\mathbf{x}_i, \mathbf{W}), y_i)$.

Proof. If $AAE(f(\mathbf{x}_i, \mathbf{W}), y_i)$ is convex, it is quasi-convex. The log function is monotonically increasing, so the composition $\log AAE(f(\mathbf{x}_i, \mathbf{W}), y_i)$ is quasi-convex.²

$\log$ is a strictly monotone function and $NAAE(f(\mathbf{x}_i, \mathbf{W}), y_i)$ is quasi-convex, so it shares the same local and global optima with $AAE(f(\mathbf{x}_i, \mathbf{W}), y_i)$. $\square$

The above theorem states that the convexity region of NAAE is consistent with AAE. To interpret this statement in another perspective, the log function is a strictly monotone function. Even if AAE is not strictly convex, NAAE still shares the same local and global optima with AAE. If we define the mapping function $f : AAE \rightarrow NAAE$, it is easy to see that f is bijective and continuous. Its inverse map f^{-1} is also continuous, so that f is an open mapping. Thus, it is easy to prove that the mapping function f is a homeomorphism to preserve all the topological properties of the given space.

The above theorems state the consistent relations among NAAE, AAE and MSE. It is proven that the greater the RO index $|\lambda|$, the larger the convex region is. Intuitively, increasing $|\lambda|$ creates tunnels for a local-search minimization procedure to travel through to a good local optimum. While NAAE preserves the robustness to MSE, theorem 9 provides insights into when NAAE is optimal than MSE.

Theorem 9. *Given training samples $\{\mathbf{X}, y\} = \{(\mathbf{x}_1, y_1), (\mathbf{x}_2, y_2), \dots, (\mathbf{x}_m, y_m)\}$ and the model $f(\mathbf{x}_i, \mathbf{W})$ with parameters $\mathbf{W}$. Define the deviation $\alpha_i = f(\mathbf{x}_i, \mathbf{W}) - y_i$, when $\alpha_i \rightarrow 1$ and $\lambda \leq -1$, both $AAE(f(\mathbf{x}_i, \mathbf{W}), y_i)$ and $NAAE(f(\mathbf{x}_i, \mathbf{W}), y_i)$ always have a larger convexity region to commit a higher chance for the better local optima than MSE.*

²Because the function f defined by $f(x) = g(U(x))$ is quasi-convex if the function U is quasiconvex and the function g is increasing.

Proof. Let $h(\mathbf{W})$ denotes the Hessian matrix of MSE (Eqn. 6.1),

$$(6.10) \quad \begin{aligned} h(\mathbf{W}) &= \frac{2}{m} \sum_{i=1}^m \left\{ \alpha_i(\mathbf{W})^2 \frac{\partial f(\mathbf{x}_i, \mathbf{W})^2}{\partial \mathbf{W}} \right. \\ &\quad \left. + \alpha_i(\mathbf{W}) \frac{\partial f^2(\mathbf{x}_i, \mathbf{W})}{\partial \mathbf{W}^2} \right\} \end{aligned}$$

Since $\lambda \leq -1$, let $diag_{eig}$ denote the diagonal matrix of the eigenvalues from SVD decomposition. $\succeq$ here means 'element-wise greater'. When $A \succeq B$, each element in A is greater than B. Then we have

$$\begin{aligned} &diag_{eig}[H(\mathbf{W})] \\ &\succeq \frac{2}{m} \left\{ (2\lambda^2 \alpha_i^2 + \lambda) \frac{\partial f(\mathbf{x}_i, \mathbf{W})^2}{\partial \mathbf{W}} - \lambda \alpha_i \frac{f^2(\mathbf{x}_i, \mathbf{W})}{\partial \mathbf{W}^2} \right\} \\ &\succeq \alpha_i^2 \frac{f(\mathbf{x}_i, \mathbf{W})^2}{\partial \mathbf{W}} + \alpha_i \frac{f^2(\mathbf{x}_i, \mathbf{W})}{\partial \mathbf{W}^2} \\ &\succeq diag_{eig}[h(\mathbf{W})] \end{aligned}$$

Briefly, when the standard deviation is large (approaches 1) and $\lambda \leq -1$, $AAE(f(\mathbf{x}_i, \mathbf{W}), y_i)$ always has larger convexity regions than MSE to better enable escape of local minima. Because $NAAE(f(\mathbf{x}_i, \mathbf{W}), y_i)$ is quasi-convex, sharing the same local and global optimum with $AAE(f(\mathbf{x}_i, \mathbf{W}), y_i)$, the above conclusions are still valid. $\square$

The expansion of the convexity region in the Hessian matrix enables NAAE to find a better local optima than MSE. This property guarantees higher probability to escape poor local optima. In the worst case, NAAE will perform as good as MSE if the convexity region shrinks as the RO index λ increase to approach 0 or the local search deviates from the "tunnel" of convex regions.

6.3.3 Control Robustness and Optimality

NAAE performs robustly while preserving the optimality by the expansion of its convex region in its Hessian matrix, the RO index λ controls both the robustness and optimality of the estimator simultaneously. When λ is a large negative number, NAAE works like a minimin operator to avoid large deviations incurred by the outliers. Meanwhile, its Hessian matrix has a larger convex region to facilitate seeking a better optimum. If λ approaches 0, a compensation from the exponential sum of the squared errors offsets the impact of the minimin operator to further focus on the larger overview. This will help to address the small deviation from other noise. NAAE performs more like MSE to grant a smooth and stable optimization procedure.

6.4 Update Strategies

As λ controls the degree of both robustness and optimality, the initialization and update strategy impacts the learning performance. Generally, large λ is preferable at the start. Different update strategies lead to three major learning methods.

6.4.1 Fixed- λ

Using a fixed large λ to optimize the exponential estimator is reported in (Lo, Gui, & Peng 2012). For NAAE, a fixed λ consistently controls the estimator at a specific level of the offset towards the minimin operator. The weights are updated using gradient descent. It is simple and work well to approximate the lower dimensional functions with outliers. When attacking the learning tasks in higher dimensions such as visual recognition, the large plateau and unstable learning procedure can lead to a failure in training (Wang, Oates, & Lo 2015; Gui, Lo, & Peng 2014).

6.4.2 Gradual Deconvexification

(Lo, Gui, & Peng 2013) proposed the Gradual Deconvexification (GDC) approach to alleviate the difficulty in finding a good value of λ for training. GDC starts with a very large λ while recording the values of the objective function before and after a pre-selected number of training epochs. If the performance converges, GDC flags the current training condition as stagnant and performs a deconvexification by reducing the current λ with a percentage r . After that, the training continues and repeats the deconvexification if necessary until a satisfied training error is achieved. If $|\lambda|$ keeps coming down to below 1, the training will set $\lambda = 0$, which actually converts to the training with MSE. All other weights are updated using gradient descent.

GDC works in a similar way as the decreasing-learning-rate approach. Training with GDC avoids fixing the initial λ during training, as the deconvexification gradually decreases λ to avoid bad optima and vulnerable extreme estimator and eventually obtain a reasonable optimum while achieving robustness. As reported in (Gui, Lo, & Peng 2014), GDC is much slower than using MSE alone. The update strategy relies on the per-defined convergence thresholds and dropping rate.

6.4.3 Adaptive Training

(Wang, Oates, & Lo 2015) proposed a novel learning method to training with RO index λ , called the Adaptive Normalized Risk-Avering Training approach (adaptive training). Instead of manually tuning λ like GDC, they learn λ adaptively through error backpropagation by considering λ as a parameter instead of a hyperparameter. The learning procedure is standard batch gradient descent.

$$(6.11) \quad \frac{dl(W, \lambda)}{d\lambda} \approx \frac{q}{\lambda} (L_P\text{-norm error} - E)$$

As shown in the above equation, the gradient on λ is approximately the difference between the exponentialized estimator E (e.g., NAAE in this paper) and the standard L_p -norm error. Considering $\lambda < 0$ as for NAAE. When E is larger, λ decreases to enlarge the convexity region, facilitating the search in the error space for better optima. When E is smaller, the learned parameters are seemingly going through the optimal "tunnel" for better optima. λ then increases and helps the weights not deviate far from the manifold of the standard L_p -norm error to make the error space stable without large plateaus. We note that adaptive training also adjusts λ to control the robustness level in a similarly 'smart' way. If E is larger than the L_p -norm error, λ decreases to further ignore more samples with large deviations to let the learning function gain a broader overview rather than focusing on the outliers. When E is small, outliers with large deviation are almost eliminated. λ increases to approach 0, making the learning procedure more accurate and stable by considering the effect of more training samples rather than the extreme value. The adaptive training approach has more flexibility. It keeps searching the error space near the manifold of the L_p -norm error to find better optima in a way of competing with, and at the same time relying on, the standard L_p -norm error space.

We consider the fixed λ approach and adaptive training as the update strategies for the RO index λ in our experiments. For Adaptive training, the final loss function is

$$(6.12) \quad l(W, \lambda) = \frac{1}{\lambda} \log \frac{1}{m} \sum_{i=1}^m e^{\lambda(f(\mathbf{x}_i, \mathbf{W}) - y_i)^2} + a|\lambda|^{-1}$$

The penalty weight a controls the convergence speed by penalizing small $|\lambda|$. Smaller

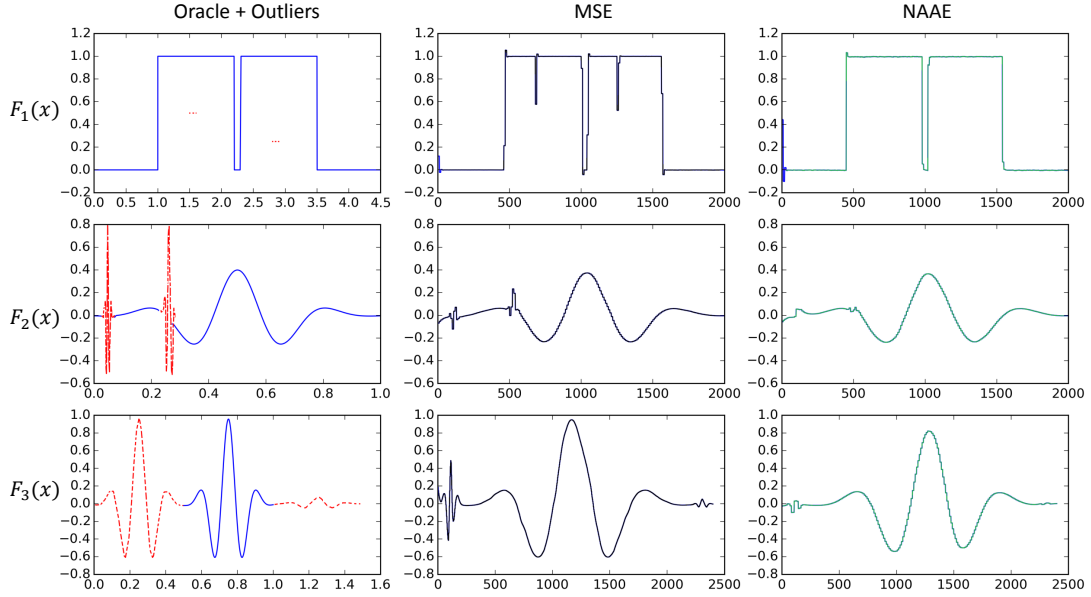


FIG. 6.1. Function approximation using MSE and NAAE on the training set. Three columns are (a). oracle function with outliers (L), (b). approximation results using MSE (M) and (c). NAAE (R). In column (a), the blue lines plot the oracle function, the red dots plot the outliers.

a emphasizes tuning λ to allow faster convergence speed between NAAE and MSE. Larger a forces larger $|\lambda|$ for a better chance to find a better optimum while skipping the outliers, but runs the risk of plateaus and deviating far from the stable error space.

6.5 Experiments and Analysis

As an error estimator, we test NAAE on the function approximation and digits recognition tasks with two different neural networks architectures, multiple layer perceptrons (MLPs) and convolutional networks (ConvNets), respectively. We limit all the experiments in very simple settings, using gradient descent with batch gradient descent and weight de-

cay or dropout (Srivastava *et al.* 2014).

6.5.1 Function Approximation

For function approximation tasks, we design three typical non-convex functions with different types of outliers to test the robustness and optimality achieved by NAAE.

1). The notch function is defined by

$$f_1(x) = \begin{cases} 0 & \text{if } x \in [0, 1.0] \cup [2.2, 2.3] \cup [3.5, 4.5] \\ 1 & \text{otherwise} \end{cases}$$

The outliers in the middle range are generated by the function

$$f_1^*(x) = \begin{cases} 0.25 & \text{if } x \in [2.8, 2.81] \\ 0.5 & \text{if } x \in [1.5, 1.51] \end{cases}$$

where $x \in X = [0, 4.5]$. x_i are obtained by random sampling 2000 non-repeatable numbers from X with a uniform distribution, and the corresponding output values y_k are computed by $f(x)$ and $f^*(x)$. The overall function is $F_1(x) = f_1(x) + f_1^*(x)$. The training data with 2000 (x_i, y_i) pairs is chosen to perform the notch function approximation $f(x)$ with the corruption of the outliers $f^*(x)$. MLPs with 1:16:1 architectures are initiated to all training sessions.

2). The smooth function is defined by

$$f_2(x) = g(x, \frac{1}{6}, \frac{1}{2}, \frac{1}{6})$$

The outliers with very large derivations given by

$$f_2^*(x) = g(x, \frac{1}{64}, \frac{1}{4}, \frac{1}{128}) + g(x, \frac{1}{64}, \frac{1}{20}, \frac{1}{128})$$

Where $x \in X = [0, 1]$, g is defined as

$$(6.13) \quad g(x) = \frac{\alpha}{\sqrt{2\pi\tau}} \cos\left(\frac{(x-\mu)\pi}{\tau}\right) e^{-\frac{(x-\mu)^2}{2\tau^2}}$$

The overall function is $F_2(x) = f_2(x) + f_2^*(x)$. The input values x_i are selected by sampling 2000 numbers from a uniformly distributed grid on X . The corresponding output values y_i are computed by $f(x)$ and $f^*(x)$. The training data with 2000 (x_i, y_i) pairs is chosen to perform the smooth function approximation with two fine-feature intrusions. MLPs with 1:15:1 architectures are applied.

3). Define a smooth function by

$$F_3(x) = g(x, \frac{1}{5}, \frac{1}{4}, \frac{1}{12}) + g(x, \frac{1}{5}, \frac{3}{4}, \frac{1}{12}) + g(x, \frac{1}{64}, \frac{5}{4}, \frac{1}{12})$$

The outliers from the under-sampled points are given by the lower sampling frequency. $x \in X = [0, 1.5]$, the outliers x_i^* are collected by sampling 200 numbers from a uniform distributed grid on $[0, 0.5]$ and 200 numbers from a uniform distributed grid on $[1.0, 1.5]$. The smooth function is generated by 2000 numbers of points from a uniform distributed grid on $(0.5, 1.0)$. So the oracle function is $f_3(x) = F_3(x), x \in (0.5, 1)$ and the outlier generating function is $f_3(x) = F_3(x), x \in [0, 0.5] \cup [1, 1.5]$. The output y_i are computed by $F_3(x)$. The training data with 2400 (x_i, y_i) pairs is the function contaminated by the unevenly under-sampled segments. MLPs with 1:12:1 architecture are initiated to all training sessions.

The sample functions represent three typical sources of outliers (Figure 6.1 (a)). The

Table 6.1. Training and test MSE of approximation on the three sample functions. For NAAE, λ is updated by the fixed- λ and adaptive learning strategies. The average performance over 10 runs are reported.

	MSE		NAAE-fixed- λ		NAAE-adaptive learning	
	Training	Test	Training	Test	Training	Test
$F_1(x)$	0.02	0.022	0.021	0.019	0.021	0.018
$F_2(x)$	0.013	0.029	0.005	0.005	0.005	0.005
$F_3(x)$	0.012	0.037	0.045	0.01	0.045	0.01

outliers in $F_1(x)$ are in the range of the uncontaminated data but appearing at wrong positions. This sometimes occurs when the labels are normal but maybe corrupted or inaccurate. $F_2(x)$ is affected by outliers with significant large deviations (even exceeding the bounds), indicating the situations where the labels contain some unknown samples. $F_3(x)$ includes seemingly uncontaminated but shifted and under-sampled data. This may happen among heterogeneous observations from several experiments where the samples are highly sparse and biased, and the main function needs to get rid of their influence.

In the following experiments, we use both fixed λ and adaptive learning methods to update the RO index. λ starts at -10^3 . NAAE is optimized with batch gradient descent. Batch size is fixed at 20. For each of the functions, we generate the test set consisting 2000 samples from the oracle function $f_i(x)$, ($i \in 1, 2, 3$) respectively.

The approximation results of 10 runs are summarized in Table 6.1. NAAE is more robust to these three types of outliers with both update strategies. We did not find significant differences between fixed- λ and adaptive learning on these experiments. Training with MSE always achieves lower training error but with a higher test error, which is also known as overfitting. NAAE eliminates the impact of outliers efficiently to achieve better test errors. It is interesting that when approximating the smooth function with fine-features ($F_2(x)$), NAAE achieves better training and test error simultaneously. This perhaps due to both the robustness and optimality of NAAE. Optimizing MSE is unable to get rid of the

outliers. The high non-convexity of the compound function also leads to the problem of local optimum. While ignoring the outliers, NAAE seeks the better local optima due to the expansion of the convex region to obtain better performance. The empirical breakdown point of NAAE on these three samples are 0.4%, 8.4% and 16.7%. The training results are shown in Figure 6.1.

6.5.2 Visual Digits Recognition

We further investigate NAAE in a typical learning problem, the digits recognition tasks on the MNIST dataset. The MNIST dataset (LeCun *et al.* 1998) consists of hand written digits 0-9 which are 28x28 in size. There are 60,000 training images and 10,000 testing images in total. We use 10,000 images in the training set for validation to select the hyperparameters and report the performance on the test set. We test our method on this dataset without data augmentation.

The NAAE function is minimized by batch gradient descent with momentum at 0.9. The learning rate and l_2 penalty are fixed at 0.5 and 0.05. The penalty weight α are selected in $\{1, 0.1, 0.001\}$ on validation sets respectively. The initial λ is fixed at -100. We use the hold-out validation set to select the best model, which is used to make predictions on the test set. All experiments are implemented quite easily in Python and Theano to obtain GPU acceleration (Bastien *et al.* 2012).

On the MNIST dataset we use the same structure of LeNet5 with two convolutional max-pooling layers but followed by only one fully connected layer and a densely connected softmax layer. The first convolutional layer has 20 feature maps of size 5×5 and max-pooled by 2×2 non-overlapping windows. The second convolutional layer has 50 feature maps with the same convolutional and max-pooling size. The fully connected layer has

³(1)(Mairal *et al.* 2014);(2)(Lee *et al.* 2014); (3)(Zeiler & Fergus 2013);(4)(Jarrett *et al.* 2009)

Method ³	Error %
Convolutional Kernel Networks ⁽¹⁾	0.39
Deeply Supervised Nets + dropout ⁽²⁾	0.39
ConvNets + dropout ⁽³⁾	0.55
large ConvNets, unsup pretraining ⁽⁴⁾	0.53
ConvNets + NAAE (Ours)	0.53
ConvNets + NAAE + dropout (Ours)	0.39

Table 6.2. Test set misclassification rates of the best methods that utilized convolutional networks on the original MNIST dataset using single model.

500 hidden units. An l_2 prior was used with the strength 0.05 in the Softmax layer. Trained by ANRAT, we can obtain a test set error of 0.53%, which is the best result we are aware of that does not use dropout on the pure ConvNets. With dropout, our method achieve the same performance with the state-of-art at 0.39% error. We summarize the best published results on the standard MNIST dataset in Table 6.2.

The above results demonstrate both the optimality and robustness of NAAE. The improvements in MNIST are generally due to better regularization techniques with good optimization strategy. With simple experiment settings and network architectures, NAAE enables the simple ConvNets to optimally learn the models while getting rid of specific outlier samples to enhance the generalization capability.

To better evaluate the robustness of NAAE, we used a fraction of a permutation of the labels to add outliers among the samples but without fixed patterns like (Reed *et al.* 2014). The learning model is the same as the last experiment on the MNIST dataset. The noise fraction ranges from 0% to 35%. Figure 6.2 shows that our NAAE with the adaptive training method provides a significant benefit in the case of permuted labels. The consistently lower fitting MSE also indicates its optimality. As the noise level increases up to 30%, NAAE provides the benefit in this high-noise regime, but is only slightly better than or

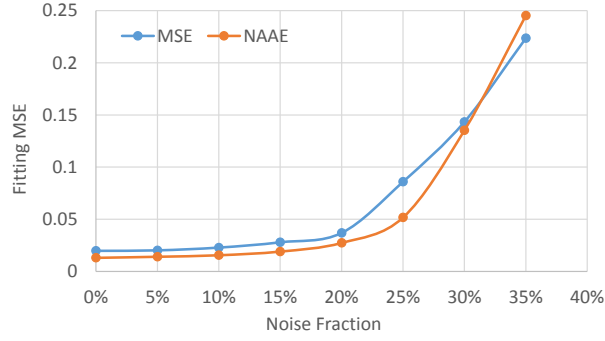


FIG. 6.2. Digit recognition error rates versus percent corrupted labels.

same with training using MSE overall. When the noise fraction is larger than 30%, NAAE perform even worse than MSE due to the failure of robustness to identify the outliers, while the convex region still expands to 'optimally' fit the noisy model, thus incurs overfitting.

6.6 Conclusions

We introduced the Normalized Anomaly-Averting Estimator by pushing the robust-optimal (RO) index λ to $-\infty$. It is robust to outliers due to its quasi-minimin functionality. The robustness is realized and controlled by its adaptive RO index without any predefined threshold. Its optimality is guaranteed by the expansion of the convexity region in its Hessian matrix to largely avoid the local optima. We provide both theoretical and empirical support for its robustness and optimality.

In future work, it may be promising to consider updating λ between $-\infty$ to ∞ to achieve robustness to both small noise and large outliers while preserving optimality, provide the risk and population risk bounds and extend our approach to other areas like aircraft and robotics control. It is also interesting to augment large-scale training for detection (e.g. ILSVRC and speech) with unlabeled and more weakly-labeled images/corpus.

Chapter 7

CONCLUSION AND FUTURE WORK

We introduced three different perspectives/frameworks to model temporal data (primarily time series data) and learn the representations. Motivated by the internal correlation embedded in time series and intrinsic property of BoP representations, we proposed time warping SAX to integrate the temporal correlation when building SAX words and BoP representations. Pooling SAX-BoP with Boosting approach suppose to solve classification problems on the multivariate vital signs time series. Instead of majority voting, the Boosting algorithm is applied to significantly improve the performance. When the data has strong internal temporal correlations, time-warping always tends to work better than the vanilla SAX. Pooling SAX-BoP with Boosting approach is motivated to solve the problem of representation modeling on multivariate time series. It demonstrates its effectiveness and efficiency compared with the current state-of-the-art approaches especially on the multivariate physiological data. Instead of majority voting, the Boosting algorithm is applied to significantly improve the performance. When the number of channel is smaller than 6, Pooling SAX-BoP is much worth to try.

Imaging time series is an off-line approach for spatially encoding the temporal patterns for classification with deep learning frameworks. We created a pipeline for converting trajectory and time series data into novel representations, GAF and MTF images,

and extracted high-level features from these using ConvNets. The features were subsequently used for classification. We demonstrated that our approach yields competitive results when compared to state-of-the-art methods by searching a relatively small parameter space. We found that GAF-MTF multi-channel images are scalable to larger numbers of quasi-orthogonal features that yield more comprehensive images. Our analysis of high-level features learned from ConvNets suggested Tiled ConvNets work like multi-frequency moving averages that benefit from the 2D temporal dependency that is preserved by the Gramian matrix. When the sample length is standard but the sample size is large, imaging + deep learning demonstrates the supreme performance on a variety type of temporal data.

Modeling and feature learning using deep neural networks are widely explored recently. We introduce a set of novel error estimators with learning methods, NRAE and ANRAT to help train deep neural networks. Theoretically, we prove the effectiveness of Normalized Risk-Averting Error on its arithmetic bound, global convexity and local convexity lower-bounded by standard L_p -norm error when convexity index $\lambda \geq 1$. By analyzing the gradient on λ , we explained the reason why using back propagation on λ works. The experiments on deep/shallow network layouts demonstrate comparable or better performance with the same experimental settings among pure ConvNets and MLP + batch SGD on MSE and CE (with or without dropout). Other than unsupervised pretraining, it provides a new perspective to address the non-convex optimization strategy in DNNs.

NAAE is a variation of NRAE with the constraint $\lambda < 0$. That is, we push the robust-optimal (RO) index λ to $-\infty$. It is robust to outliers due to its quasi-minimin functionality. The robustness is realized and controlled by its adaptive RO index without any predefined threshold. Its optimality is guaranteed by the expansion of the convexity region in its Hessian matrix to largely avoid the local optima. We provide both theoretical and empirical support for its robustness and optimality. NRAE/NAAE with ANRAT are good alternatives to MSE/cross entropy in fitting and learning problems, especially for modeling temporal

data using neural nets in the end-to-end manner.

To summarize, symbolic approximation approaches discretize the temporal data as symbolic representation, facilitating to extract the bag-of-features for modeling temporal correlation. By symbolic approximation, a bunch of learning algorithm in natural language processing (NLP) can be adapted to temporal data mining. Imaging + deep learning framework enables the model to learn more complicated temporal dynamics. It bridges the research in computer vision and time series analysis, enabling the modeling and representation learning on temporal data to benefit from the rapid development of deep learning based computer vision approaches. NRAE/NAAE and ANRAT provide another perspective to address the optimality and robustness with a unified framework for the non-convex optimization problem in deep learning.

As for the future work, we will explore the idea of imaging and symbolization for representation learning on the multivariate temporal data. How to learn the representations which will benefit temporal reasoning/inference is also a key topic. For the Non-convex optimization problem, more effort is worth to be spent on much large data set (e.g. ImageNet) with more complex vision models (e.g. GooleNet, VGG-Net) and recurrent models (e.g. Deep LSTM, Deep GRU) to better understand its potential on static and temporal data learning. Its robustness is also crucial to control theory.

REFERENCES

- [1] Abdel-Hamid, O.; Mohamed, A.-r.; Jiang, H.; and Penn, G. 2012. Applying convolutional neural networks concepts to hybrid nn-hmm model for speech recognition. In *Acoustics, Speech and Signal Processing (ICASSP), 2012 IEEE International Conference on*, 4277–4280. IEEE.
- [2] Abdel-Hamid, O.; Deng, L.; and Yu, D. 2013. Exploring convolutional neural network structures and optimization techniques for speech recognition. In *INTERSPEECH*, 3366–3370.
- [3] Aggarwal, C. C., and Yu, P. S. 2001. Outlier detection for high dimensional data. In *ACM Sigmod Record*, volume 30, 37–46. ACM.
- [4] Agrawal, R.; Faloutsos, C.; and Swami, A. 1993. *Efficient similarity search in sequence databases*. Springer.
- [5] An, J.; Chen, H.; Furuse, K.; Ohbo, N.; and Keogh, E. 2003. Grid-based indexing for large time series databases. In *Intelligent Data Engineering and Automated Learning*. Springer. 614–621.
- [6] Angiulli, F., and Pizzuti, C. 2005. Outlier mining in large high-dimensional data sets. *Knowledge and Data Engineering, IEEE Transactions on* 17(2):203–215.
- [7] Antunes, C. M., and Oliveira, A. L. 2001. Temporal data mining: An overview. In *KDD Workshop on Temporal Data Mining*, 1–13.
- [8] Athitsos, V., and Sclaroff, S. 2005. Boosting nearest neighbor classifiers for multi-class recognition. In *Computer Vision and Pattern Recognition-Workshops, 2005. CVPR Workshops. IEEE Computer Society Conference on*, 45–45. IEEE.

- [9] Bakshi, B. R., and Stephanopoulos, G. 1995. Reasoning in time: Modeling, analysis, and pattern recognition of temporal process trends. *Advances in Chemical Engineering* 22:485–548.
- [10] Bandt, C., and Pompe, B. 2002. Permutation entropy: a natural complexity measure for time series. *Physical Review Letters* 88(17):174102.
- [11] Bankó, Z., and Abonyi, J. 2012. Correlation based dynamic time warping of multivariate time series. *Expert Systems with Applications* 39(17):12814–12823.
- [12] Bastien, F.; Lamblin, P.; Pascanu, R.; Bergstra, J.; Goodfellow, I. J.; Bergeron, A.; Bouchard, N.; and Bengio, Y. 2012. Theano: new features and speed improvements. Deep Learning and Unsupervised Feature Learning NIPS 2012 Workshop.
- [13] Baydogan, M. G., and Runger, G. 2014. Learning a symbolic representation for multivariate time series classification. *Data Mining and Knowledge Discovery* 1–23.
- [14] Baydogan, M. G.; Runger, G.; and Tuv, E. 2013. A bag-of-features framework to classify time series. *Pattern Analysis and Machine Intelligence, IEEE Transactions on* 35(11):2796–2802.
- [15] Bengio, Y.; Lamblin, P.; Popovici, D.; Larochelle, H.; et al. 2007. Greedy layer-wise training of deep networks. *Advances in neural information processing systems* 19:153.
- [16] Bengio, Y.; Yao, L.; Alain, G.; and Vincent, P. 2013. Generalized denoising auto-encoders as generative models. In *Advances in Neural Information Processing Systems*, 899–907.
- [17] Blake, A., and Zisserman, A. 1987. *Visual reconstruction*, volume 2. MIT press Cambridge.

- [18] Box, G. E.; Jenkins, G. M.; and Reinsel, G. C. 2013. *Time series analysis: forecasting and control*. John Wiley & Sons.
- [19] Camerra, A.; Palpanas, T.; Shieh, J.; and Keogh, E. 2010. isax 2.0: Indexing and mining one billion time series. In *Data Mining (ICDM), 2010 IEEE 10th International Conference on*, 58–67. IEEE.
- [20] Campanharo, A. S.; Sirer, M. I.; Malmgren, R. D.; Ramos, F. M.; and Amaral, L. A. N. 2011. Duality between time series and networks. *PloS one* 6(8):e23378.
- [21] Chan, F.-P.; Fu, A.-C.; and Yu, C. 2003. Haar wavelets for efficient similarity search of time-series: with and without time warping. *Knowledge and Data Engineering, IEEE Transactions on* 15(3):686–705.
- [22] Chen, D. S., and Jain, R. C. 1994. A robust backpropagation learning algorithm for function approximation. *Neural Networks, IEEE Transactions on* 5(3):467–479.
- [23] Choromanska, A.; Henaff, M.; Mathieu, M.; Arous, G. B.; and LeCun, Y. 2014. The loss surface of multilayer networks. *arXiv preprint arXiv:1412.0233*.
- [24] Clark, S. P. 1990. Estimating the fractal dimension of chaotic time series. *Lincoln Laboratory Journal* 3(1).
- [25] Coates, A., and Ng, A. Y. 2011. Selecting receptive fields in deep networks. In *Advances in Neural Information Processing Systems*, 2528–2536.
- [26] Cohen, N., and Shashua, A. 2014. Simnets: A generalization of convolutional networks. *arXiv preprint arXiv:1410.0781*.
- [27] Dauphin, Y. N.; Pascanu, R.; Gulcehre, C.; Cho, K.; Ganguli, S.; and Bengio, Y. 2014. Identifying and attacking the saddle point problem in high-dimensional non-

- convex optimization. In *Advances in Neural Information Processing Systems*, 2933–2941.
- [28] Daw, C. S.; Finney, C. E. A.; and Tracy, E. R. 2003. A review of symbolic analysis of experimental data. *Review of Scientific Instruments* 74(2):915–930.
- [29] Deng, L.; Abdel-Hamid, O.; and Yu, D. 2013. A deep convolutional neural network using heterogeneous pooling for trading acoustic invariance with phonetic confusion. In *Acoustics, Speech and Signal Processing (ICASSP), 2013 IEEE International Conference on*, 6669–6673. IEEE.
- [30] Deng, L.; Li, J.; Huang, J.-T.; Yao, K.; Yu, D.; Seide, F.; Seltzer, M.; Zweig, G.; He, X.; Williams, J.; et al. 2013. Recent advances in deep learning for speech research at microsoft. In *Acoustics, Speech and Signal Processing (ICASSP), 2013 IEEE International Conference on*, 8604–8608. IEEE.
- [31] Deng, L.; Hinton, G.; and Kingsbury, B. 2013. New types of deep neural network learning for speech recognition and related applications: An overview. In *Acoustics, Speech and Signal Processing (ICASSP), 2013 IEEE International Conference on*, 8599–8603. IEEE.
- [32] Ding, H.; Trajcevski, G.; and Scheuermann, P. 2008. Efficient maintenance of continuous queries for trajectories. *GeoInformatica* 12(3):255–288.
- [33] Donner, R. V.; Zou, Y.; Donges, J. F.; Marwan, N.; and Kurths, J. 2010. Recurrence networks: a novel paradigm for nonlinear time series analysis. *New Journal of Physics* 12(3):033025.
- [34] Donner, R. V.; Small, M.; Donges, J. F.; Marwan, N.; Zou, Y.; Xiang, R.; and Kurths,

- J. 2011. Recurrence-based time series analysis by means of complex network methods. *International Journal of Bifurcation and Chaos* 21(04):1019–1046.
- [35] Erhan, D.; Bengio, Y.; Courville, A.; Manzagol, P.-A.; Vincent, P.; and Bengio, S. 2010. Why does unsupervised pre-training help deep learning? *The Journal of Machine Learning Research* 11:625–660.
- [36] Faloutsos, C.; Ranganathan, M.; and Manolopoulos, Y. 1994. *Fast subsequence matching in time-series databases*, volume 23. ACM.
- [37] Fan, R.-E.; Chang, K.-W.; Hsieh, C.-J.; Wang, X.-R.; and Lin, C.-J. 2008. Liblinear: A library for large linear classification. *The Journal of Machine Learning Research* 9:1871–1874.
- [38] Feng, J.; Xu, H.; Mannor, S.; and Yan, S. 2014. Robust logistic regression and classification. In *Advances in Neural Information Processing Systems*, 253–261.
- [39] Flash, T., and Hochner, B. 2005. Motor primitives in vertebrates and invertebrates. *Current opinion in neurobiology* 15(6):660–666.
- [40] Freund, Y., and Schapire, R. E. 1995. A decision-theoretic generalization of on-line learning and an application to boosting. In *Computational learning theory*, 23–37. Springer.
- [41] Freund, Y.; Schapire, R. E.; et al. 1996. Experiments with a new boosting algorithm. In *ICML*, volume 96, 148–156.
- [42] Gan, Z.; Henao, R.; Carlson, D.; and Carin, L. 2015. Learning deep sigmoid belief networks with data augmentation. In *Proceedings of the Eighteenth International Conference on Artificial Intelligence and Statistics*, 268–276.

- [43] Gandhi, S. R., and Oates, T. 2015. A generative model for time series based on multiple normal distributions. *Master thesis*.
- [44] Goodfellow, I. J.; Warde-Farley, D.; Lamblin, P.; Dumoulin, V.; Mirza, M.; Pascanu, R.; Bergstra, J.; Bastien, F.; and Bengio, Y. 2013a. Pylearn2: a machine learning research library. *arXiv preprint arXiv:1308.4214*.
- [45] Goodfellow, I. J.; Warde-Farley, D.; Mirza, M.; Courville, A.; and Bengio, Y. 2013b. Maxout networks. *arXiv preprint arXiv:1302.4389*.
- [46] Goodfellow, I.; Pouget-Abadie, J.; Mirza, M.; Xu, B.; Warde-Farley, D.; Ozair, S.; Courville, A.; and Bengio, Y. 2014. Generative adversarial nets. In *Advances in Neural Information Processing Systems*, 2672–2680.
- [47] Gui, Y.; Lo, J. T.-H.; and Peng, Y. 2014. A pairwise algorithm for training multi-layer perceptrons with the normalized risk-averting error criterion. In *Neural Networks (IJCNN), 2014 International Joint Conference on*, 358–365. IEEE.
- [48] Hannun, A.; Case, C.; Casper, J.; Catanzaro, B.; Diamos, G.; Elsen, E.; Prenger, R.; Satheesh, S.; Sengupta, S.; Coates, A.; et al. 2014. Deepspeech: Scaling up end-to-end speech recognition. *arXiv preprint arXiv:1412.5567*.
- [49] He, K.; Zhang, X.; Ren, S.; and Sun, J. 2015. Delving deep into rectifiers: Surpassing human-level performance on imagenet classification. *arXiv preprint arXiv:1502.01852*.
- [50] Hermansky, H. 1990. Perceptual linear predictive (plp) analysis of speech. *the Journal of the Acoustical Society of America* 87(4):1738–1752.
- [51] Hilbert, D. 1891. Ueber die stetige abbildung einer line auf ein flächenstück. *Mathematische Annalen* 38(3):459–460.

- [52] Hinton, G.; Deng, L.; Yu, D.; Dahl, G. E.; Mohamed, A.-r.; Jaitly, N.; Senior, A.; Vanhoucke, V.; Nguyen, P.; Sainath, T. N.; et al. 2012. Deep neural networks for acoustic modeling in speech recognition: The shared views of four research groups. *Signal Processing Magazine, IEEE* 29(6):82–97.
- [53] Hinton, G.; Osindero, S.; and Teh, Y.-W. 2006. A fast learning algorithm for deep belief nets. *Neural computation* 18(7):1527–1554.
- [54] Hornik, K.; Stinchcombe, M.; and White, H. 1989. Multilayer feedforward networks are universal approximators. *Neural networks* 2(5):359–366.
- [55] Hubel, D. H., and Wiesel, T. N. 1962. Receptive fields, binocular interaction and functional architecture in the cat’s visual cortex. *The Journal of physiology* 160(1):106.
- [56] Jacobson, D. H. 1973. Optimal stochastic linear systems with exponential performance criteria and their relation to deterministic differential games. *Automatic Control, IEEE Transactions on* 18(2):124–131.
- [57] Jaderberg, M.; Simonyan, K.; Zisserman, A.; and Kavukcuoglu, K. 2015. Spatial transformer networks. *arXiv preprint arXiv:1506.02025*.
- [58] Jarrett, K.; Kavukcuoglu, K.; Ranzato, M.; and LeCun, Y. 2009. What is the best multi-stage architecture for object recognition? In *Computer Vision, 2009 IEEE 12th International Conference on*, 2146–2153. IEEE.
- [59] Kalpakis, K.; Gada, D.; and Puttagunta, V. 2001. Distance measures for effective clustering of arima time-series. In *Data Mining, 2001. ICDM 2001, Proceedings IEEE International Conference on*, 273–280. IEEE.
- [60] Karpathy, A., and Fei-Fei, L. 2015. Deep visual-semantic alignments for generating

- image descriptions. In *The IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*.
- [61] Kavukcuoglu, K.; Sermanet, P.; Boureau, Y.-L.; Gregor, K.; Mathieu, M.; and Cun, Y. L. 2010. Learning convolutional feature hierarchies for visual recognition. In *Advances in neural information processing systems*, 1090–1098.
- [62] Keogh, E. J., and Pazzani, M. J. 1998. An enhanced representation of time series which allows fast and accurate classification, clustering and relevance feedback. In *KDD*, volume 98, 239–243.
- [63] Keogh, E. J., and Pazzani, M. J. 2000. Scaling up dynamic time warping for datamining applications. In *Proceedings of the sixth ACM SIGKDD international conference on Knowledge discovery and data mining*, 285–289. ACM.
- [64] Keogh, E.; Chakrabarti, K.; Pazzani, M.; and Mehrotra, S. 2001. Dimensionality reduction for fast similarity search in large time series databases. *Knowledge and Information Systems* 3(3):263–286.
- [65] Keogh, E.; Xi, X.; Wei, L.; and Ratanamahatana, C. A. 2006. The ucr time series classification/clustering homepage. URL= http://www.cs.ucr.edu/~eamonn/time_series_data.
- [66] Keogh, E.; Xi, X.; Wei, L.; and Ratanamahatana, C. A. 2011. The ucr time series classification/clustering homepage. URL= http://www.cs.ucr.edu/~eamonn/time_series_data.
- [67] Keogh, E.; Lonardi, S.; and Ratanamahatana, C. A. 2004. Towards parameter-free data mining. In *Proceedings of the tenth ACM SIGKDD international conference on Knowledge discovery and data mining*, 206–215. ACM.

- [68] Kononenko, I. 2001. Machine learning for medical diagnosis: history, state of the art and perspective. *Artificial Intelligence in medicine* 23(1):89–109.
- [69] Korn, F.; Jagadish, H. V.; and Faloutsos, C. 1997. Efficiently supporting ad hoc queries in large datasets of time sequences. *ACM SIGMOD Record* 26(2):289–300.
- [70] Kriegel, H.-P.; Zimek, A.; et al. 2008. Angle-based outlier detection in high-dimensional data. In *Proceedings of the 14th ACM SIGKDD international conference on Knowledge discovery and data mining*, 444–452. ACM.
- [71] Krizhevsky, A., and Hinton, G. 2009. Learning multiple layers of features from tiny images. *Computer Science Department, University of Toronto, Tech. Rep* 1(4):7.
- [72] Krizhevsky, A.; Sutskever, I.; and Hinton, G. E. 2012. Imagenet classification with deep convolutional neural networks. In *Advances in neural information processing systems*, 1097–1105.
- [73] Kumar, N.; Lolla, V. N.; Keogh, E. J.; Lonardi, S.; and Ratanamahatana, C. A. 2005. Time-series bitmaps: a practical visualization tool for working with large time series databases. In *SDM*, 531–535. SIAM.
- [74] Künsch, H. R.; Stefanski, L. A.; and Carroll, R. J. 1989. Conditionally unbiased bounded-influence estimation in general regression models, with applications to generalized linear models. *Journal of the American Statistical Association* 84(406):460–466.
- [75] Längkvist, M.; Karlsson, L.; and Loutfi, A. 2014. A review of unsupervised feature learning and deep learning for time-series modeling. *Pattern Recognition Letters* 42:11–24.
- [76] Larochelle, H.; Bengio, Y.; Louradour, J.; and Lamblin, P. 2009. Exploring strategies for training deep neural networks. *The Journal of Machine Learning Research* 10:1–40.

- [77] Lawrence, S.; Giles, C. L.; Tsoi, A. C.; and Back, A. D. 1997. Face recognition: A convolutional neural-network approach. *Neural Networks, IEEE Transactions on* 8(1):98–113.
- [78] LeCun, Y., and Bengio, Y. 1995. Convolutional networks for images, speech, and time series. *The handbook of brain theory and neural networks* 3361.
- [79] LeCun, Y.; Boser, B.; Denker, J. S.; Henderson, D.; Howard, R. E.; Hubbard, W.; and Jackel, L. D. 1989. Backpropagation applied to handwritten zip code recognition. *Neural computation* 1(4):541–551.
- [80] LeCun, Y.; Bottou, L.; Bengio, Y.; and Haffner, P. 1998. Gradient-based learning applied to document recognition. *Proceedings of the IEEE* 86(11):2278–2324.
- [81] LeCun, Y.; Kavukcuoglu, K.; and Farabet, C. 2010. Convolutional networks and applications in vision. In *Circuits and Systems (ISCAS), Proceedings of 2010 IEEE International Symposium on*, 253–256. IEEE.
- [82] Lee, J.-G.; Han, J.; Li, X.; and Gonzalez, H. 2008. Traiclass: trajectory classification using hierarchical region-based and trajectory-based clustering. *Proceedings of the VLDB Endowment* 1(1):1081–1094.
- [83] Lee, C.-Y.; Xie, S.; Gallagher, P.; Zhang, Z.; and Tu, Z. 2014. Deeply-supervised nets. *arXiv preprint arXiv:1409.5185*.
- [84] Leggetter, C. J., and Woodland, P. C. 1995. Maximum likelihood linear regression for speaker adaptation of continuous density hidden markov models. *Computer Speech & Language* 9(2):171–185.
- [85] Liano, K. 1996. Robust error measure for supervised neural network learning with outliers. *Neural Networks, IEEE Transactions on* 7(1):246–250.

- [86] Lin, J.; Keogh, E.; Lonardi, S.; and Chiu, B. 2003. A symbolic representation of time series, with implications for streaming algorithms. In *Proceedings of the 8th ACM SIGMOD workshop on Research issues in data mining and knowledge discovery*, 2–11. ACM.
- [87] Lin J, Williamson S, B. K. D. D. 2012. *Pattern recognition in time series*. Chapman & Hall, To appear.
- [88] Lin, J.; Khade, R.; and Li, Y. 2012. Rotation-invariant similarity in time series using bag-of-patterns representation. *Journal of Intelligent Information Systems* 39(2):287–315.
- [89] Liu, W., and Floudas, C. A. 1993. A remark on the gop algorithm for global optimization. *Journal of Global Optimization* 3(4):519–521.
- [90] Lo, J. T., and Bassu, D. 2001. Training multilayer perceptrons in the presence of measurement outliers. In *Neural Networks, 2001. Proceedings. IJCNN'01. International Joint Conference on*, volume 3, 2030–2035. IEEE.
- [91] Lo, J. T.-H.; Gui, Y.; and Peng, Y. 2012. Overcoming the local-minimum problem in training multilayer perceptrons with the nrae training method. In *Advances in Neural Networks–ISNN 2012*. Springer. 440–447.
- [92] Lo, J. T.-H.; Gui, Y.; and Peng, Y. 2013. Overcoming the local-minimum problem in training multilayer perceptrons by gradual deconvexification. In *Neural Networks (IJCNN), The 2013 International Joint Conference on*, 1–6. IEEE.
- [93] Lo, J. T.-H. 2010. Convexification for data fitting. *Journal of global optimization* 46(2):307–315.

- [94] Lu, G.; Yang, F.; Taylor, J.; and Stein, J. 2009. A comparison of photoplethysmography and ecg recording to analyse heart rate variability in healthy subjects. *Journal of medical engineering & technology* 33(8):634–641.
- [95] Mairal, J.; Koniusz, P.; Harchaoui, Z.; and Schmid, C. 2014. Convolutional kernel networks. In *Advances in Neural Information Processing Systems*, 2627–2635.
- [96] Manly, B. F. 2006. *Randomization, bootstrap and Monte Carlo methods in biology*, volume 70. CRC Press.
- [97] Martens, J., and Grosse, R. 2015. Optimizing neural networks with kronecker-factored approximate curvature. *arXiv preprint arXiv:1503.05671*.
- [98] Martens, J. 2010. Deep learning via hessian-free optimization. In *Proceedings of the 27th International Conference on Machine Learning (ICML-10)*, 735–742.
- [99] Min Lin, Qiang Chen, S. Y. 2014. Network in network. *arXiv preprint arXiv:1312.4400v3*.
- [100] Mohamed, A.-r.; Dahl, G. E.; and Hinton, G. 2012. Acoustic modeling using deep belief networks. *Audio, Speech, and Language Processing, IEEE Transactions on* 20(1):14–22.
- [101] Moon, B.; Jagadish, H. V.; Faloutsos, C.; and Saltz, J. H. 2001. Analysis of the clustering properties of the hilbert space-filling curve. *Knowledge and Data Engineering, IEEE Transactions on* 13(1):124–141.
- [102] Mörchen, F., and Ultsch, A. 2006. Finding persisting states for knowledge discovery in time series. In *From Data and Information Analysis to Knowledge Engineering*. Springer. 278–285.

- [103] Nair, V., and Hinton, G. E. 2010. Rectified linear units improve restricted boltzmann machines. In *Proceedings of the 27th International Conference on Machine Learning (ICML-10)*, 807–814.
- [104] Nanopoulos, A.; Alcock, R.; and Manolopoulos, Y. 2001. Feature-based classification of time-series data. *International Journal of Computer Research* 10(3).
- [105] Ng, A. Y. 1997. Preventing” overfitting” of cross-validation data. In *ICML*, volume 97, 245–253.
- [106] Ng, A. 2011. Sparse autoencoder. *CS294A Lecture notes* 72.
- [107] Ngiam, J.; Chen, Z.; Chia, D.; Koh, P. W.; Le, Q. V.; and Ng, A. Y. 2010. Tiled convolutional neural networks. In *Advances in Neural Information Processing Systems*, 1279–1287.
- [108] Ngiam, J.; Coates, A.; Lahiri, A.; Prochnow, B.; Le, Q. V.; and Ng, A. Y. 2011. On optimization methods for deep learning. In *Proceedings of the 28th International Conference on Machine Learning (ICML-11)*, 265–272.
- [109] Oates, T.; Mackenzie, C. F.; Stansbury, L. G.; Aarabi, B.; Stein, D. M.; and Hu, P. F. 2012a. Predicting patient outcomes from a few hours of high resolution vital signs data. In *Machine Learning and Applications (ICMLA), 2012 11th International Conference on*, volume 2, 192–197. IEEE.
- [110] Oates, T.; Mackenzie, C. F.; Stein, D. M.; Stansbury, L. G.; Dubose, J.; Aarabi, B.; and Hu, P. F. 2012b. Exploiting representational diversity for time series classification. In *Machine Learning and Applications (ICMLA), 2012 11th International Conference on*, volume 2, 538–544. IEEE.

- [111] Ordóñez, P.; Oates, T.; and Desjardins, M. 2012. *Multivariate time-series analysis of physiological and clinical data*. University of Maryland at Baltimore County.
- [112] Panuccio, A.; Bicego, M.; and Murino, V. 2002. A hidden markov model-based approach to sequential data clustering. In *Structural, Syntactic, and Statistical Pattern Recognition*. Springer. 734–743.
- [113] Papadimitriou, S., and Yu, P. 2006. Optimal multi-scale patterns in time series streams. In *Proceedings of the 2006 ACM SIGMOD international conference on Management of data*, 647–658. ACM.
- [114] Perng, C.-S.; Wang, H.; Zhang, S. R.; and Parker, D. S. 2000. Landmarks: a new model for similarity-based pattern querying in time series databases. In *Data Engineering, 2000. Proceedings. 16th International Conference on*, 33–42. IEEE.
- [115] Poultney, C.; Chopra, S.; Cun, Y. L.; et al. 2006. Efficient learning of sparse representations with an energy-based model. In *Advances in neural information processing systems*, 1137–1144.
- [116] Prechelt, L. 1998. Automatic early stopping using cross validation: quantifying the criteria. *Neural Networks* 11(4):761–767.
- [117] Pregibon, D. 1982. Resistant fits for some commonly used logistic models with medical applications. *Biometrics* 485–498.
- [118] Pukelsheim, F. 1994. The three sigma rule. *The American Statistician* 48(2):88–91.
- [119] Rakthanmanon, T., and Keogh, E. 2013. Fast shapelets: A scalable algorithm for discovering time series shapelets. In *Proceedings of the thirteenth SIAM conference on data mining (SDM)*. SIAM.

- [120] Ranzato, M.; Huang, F. J.; Boureau, Y.-L.; and LeCun, Y. 2007. Unsupervised learning of invariant feature hierarchies with applications to object recognition. In *Computer Vision and Pattern Recognition, 2007. CVPR'07. IEEE Conference on*, 1–8. IEEE.
- [121] Ratanamahatana, C.; Keogh, E.; Bagnall, A. J.; and Lonardi, S. 2005. A novel bit level time series representation with implication of similarity search and clustering. In *Advances in Knowledge Discovery and Data Mining*. Springer. 771–777.
- [122] Ravi Kanth, K.; Agrawal, D.; and Singh, A. 1998. Dimensionality reduction for similarity searching in dynamic databases. In *ACM SIGMOD Record*, volume 27, 166–176. ACM.
- [123] Reed, S.; Lee, H.; Anguelov, D.; Szegedy, C.; Erhan, D.; and Rabinovich, A. 2014. Training deep neural networks on noisy labels with bootstrapping. *arXiv preprint arXiv:1412.6596*.
- [124] Reynolds, D. A., and Rose, R. C. 1995. Robust text-independent speaker identification using gaussian mixture speaker models. *Speech and Audio Processing, IEEE Transactions on* 3(1):72–83.
- [125] Riedl, M.; Müller, A.; and Wessel, N. 2013. Practical considerations of permutation entropy. *The European Physical Journal Special Topics* 222(2):249–262.
- [126] Rousseeuw, P. J. 1984. Least median of squares regression. *Journal of the American statistical association* 79(388):871–880.
- [127] Salakhutdinov, R., and Hinton, G. E. 2009. Deep boltzmann machines. In *International Conference on Artificial Intelligence and Statistics*, 448–455.
- [128] Schapire, R. 2013. Explaining adaboost. In Scholkopf, B.; Luo, Z.; and Vovk, V., eds., *Empirical Inference*. Springer Berlin Heidelberg. 37–52.

- [129] Sebastiani, P.; Ramoni, M.; Cohen, P.; Warwick, J.; and Davis, J. 1999. Discovering dynamics using bayesian clustering. In *Advances in intelligent data analysis*. Springer. 199–209.
- [130] Senin, P., and Malinchik, S. 2013. Sax-vsm: Interpretable time series classification using sax and vector space model. In *Data Mining (ICDM), 2013 IEEE 13th International Conference on*, 1175–1180. IEEE.
- [131] Silva, D. F.; Souza, V.; De, M.; and Batista, G. E. 2013. Time series classification using compression distance of recurrence plots. In *Data Mining (ICDM), 2013 IEEE 13th International Conference on*, 687–696. IEEE.
- [132] Speyer, J. L.; Deyst, J.; and Jacobson, D. 1974. Optimization of stochastic linear systems with additive measurement and process noise using exponential performance criteria. *Automatic Control, IEEE Transactions on* 19(4):358–366.
- [133] Sra, S.; Nowozin, S.; and Wright, S. J. 2012. *Optimization for machine learning*. Mit Press.
- [134] Srivastava, N.; Hinton, G.; Krizhevsky, A.; Sutskever, I.; and Salakhutdinov, R. 2014. Dropout: A simple way to prevent neural networks from overfitting. *The Journal of Machine Learning Research* 15(1):1929–1958.
- [135] Stefanski, L. A.; Carroll, R. J.; and Ruppert, D. 1986. Optimally bounded score functions for generalized linear models with applications to logistic regression. *Biometrika* 73(2):413–424.
- [136] Szegedy, C.; Liu, W.; Jia, Y.; Sermanet, P.; Reed, S.; Anguelov, D.; Erhan, D.; Vanhoucke, V.; and Rabinovich, A. 2015. Going deeper with convolutions.

- [137] Taylor, G. W. 2009. *Composable, distributed-state models for high-dimensional time series*. Ph.D. Dissertation, University of Toronto.
- [138] Tibshirani, J., and Manning, C. D. 2013. Robust logistic regression using shift parameters. *CoRR*.
- [139] Vincent, P.; Larochelle, H.; Bengio, Y.; and Manzagol, P.-A. 2008. Extracting and composing robust features with denoising autoencoders. In *Proceedings of the 25th international conference on Machine learning*, 1096–1103. ACM.
- [140] Vincent, P.; Larochelle, H.; Lajoie, I.; Bengio, Y.; and Manzagol, P.-A. 2010. Stacked denoising autoencoders: Learning useful representations in a deep network with a local denoising criterion. *The Journal of Machine Learning Research* 11:3371–3408.
- [141] Wang, J.; Liu, P.; She, M. F.; Nahavandi, S.; and Kouzani, A. 2013. Bag-of-words representation for biomedical time series classification. *Biomedical Signal Processing and Control* 8(6):634–644.
- [142] Wang, Z.; Oates, T.; and Lo, J. 2015. Adaptive normalized risk-averting training for deep neural networks. *arXiv preprint arXiv:1506.02690*.
- [143] Weng, X., and Shen, J. 2008. Classification of multivariate time series using locality preserving projections. *Knowledge-Based Systems* 21(7):581–587.
- [144] Whittle, P. 1990. *Risk-sensitive optimal control*. John Wiley & Son Ltd.
- [145] Xie, J., and Yan, W. 2007. Pattern-based characterization of time series. *International Journal of Information and Systems Science* 3(3):479–491.
- [146] Ye, L., and Keogh, E. 2009. Time series shapelets: a new primitive for data mining. In *Proceedings of the 15th ACM SIGKDD international conference on Knowledge discovery and data mining*, 947–956. ACM.

- [147] Yi, B.-K., and Faloutsos, C. 2000. Fast time sequence indexing for arbitrary lp norms. *VLDB*.
- [148] Yu, C.; Liu, Z.; McKenna, T.; Reisner, A. T.; and Reifman, J. 2006. A method for automatic identification of reliable heart rates calculated from ecg and ppg waveforms. *Journal of the American Medical Informatics Association* 13(3):309–320.
- [149] Zeiler, M. D., and Fergus, R. 2013. Stochastic pooling for regularization of deep convolutional neural networks. *arXiv preprint arXiv:1301.3557*.
- [150] Zheng, Y.; Liu, Q.; Chen, E.; Ge, Y.; and Zhao, J. L. 2014. Time series classification using multi-channels deep convolutional neural networks. In *Web-Age Information Management*. Springer. 298–310.
- [151] Zhou, F.; Torre, F.; and Hodgins, J. K. 2008. Aligned cluster analysis for temporal segmentation of human motion. In *Automatic Face & Gesture Recognition, 2008. FG'08. 8th IEEE International Conference on*, 1–7. IEEE.

