

Deep Neural Networks: Classifier for Variable Stars with Quely Detection Capability

BENNY T.-H. TSANG¹ AND WILLIAM C. SCHULTZ²

¹*Kavli Institute for Theoretical Physics, University of California, Santa Barbara, CA 93106, USA*

²*Department of Physics, University of California, Santa Barbara, CA 93106, USA*

(Received May 13, 2019; Revised —; Accepted —)

Submitted to MNRAS

BSR-2019-01

Current variable star classifiers are built only with the goal of producing the correct class labels, leaving little of their ultimate capability of deep neural networks unexplored. We present a periodic light curve classifier that combines a recurrent neural network autoencoder for the supervised feature extraction and a dual-purpose estimation network for supervised classification and quely detection. The estimation network optimizes a Gaussian mixture model in the reduced-dimensional feature space, where each Gaussian component corresponds to a variable class. The estimation network with a basic structure of a single hidden layer attains a cross-validation classification accuracy of $\sim 99\%$, comparable to the conventional approaches, random forest classifiers. With the addition of photometric features, the network is capable of detecting previously unseen types of variability with precision 0.90, recall 0.96, and an F_1 score of 0.93. The simultaneous training of the autoencoder and estimation network is found to be mutually beneficial, resulting in faster autoencoder convergence, and superior classification and quely detection performance. The estimation network also delivers adequate results even when optimized with pre-trained autoencoder features, suggesting that it can readily extend existing classifiers to provide added quely detection capabilities.

Keywords: binaries: eclipsing — methods: data analysis — methods: statistical — stars: general — stars: oscillations — techniques: photometric

1. INTRODUCTION

Efficient classification of the variability of astronomical objects is crucial to defining follow-up observations and analysis. With the advent of the next-generation surveys such as the *Kepler* Space Telescope (KST, Borucki et al. 2008; KST Science Collaboration et al. 2009) and the *Zoo* (Kepler Data Facility, ZTF, Bell et al. 2019), automatic pipelines are required to categorize an unprecedented amount of light curves into previously unseen variability classes. This is an challenging task as been applied to solve the classification problem efficiently. The identification of objects with quely variability properties, however, still remains a challenging task for visual inspection by human experts.

Classifiers are seldom trained using raw light curves due to their high dimensionality and large number of observations. Instead, features of the light curve data series are obtained, a process known as *feature extraction*. Features typically include the variable period (Blight 1976; Scargle 1982; Stassun & Torres 1996; Qian et al. 2002), amplitudes and ratios of different Fourier coefficients (Wu et al. 2015), and summary statistics of the light curve (e.g. standard deviation and skewness). Random forest (RF) classifiers trained on features obtained from photometric data have been the dominant method for recent classification efforts (Petro et al. 2011; Dubat et al. 2011; Blom et al. 2012; Qian et al. 2014; Alsic et al. 2014; Bailer-Jones 2016; Jayasingh et al. 2018; Poldos et al. 2018). The probabilistic RF method, which takes into account uncertainties in both features and labels (Petro et al. 2019), is also used in recent works requiring RF's re-

for us. However, manual selection of features still requires tedious human intervention.

Deep artificial neural networks have attracted recent attention in the physical sciences due to their ability to achieve meaningful data representations without human input. Astronomers also use galaxy morphology classification (Dietrich et al. 2015; Willett et al. 2017), transient and exoplanet detection (Cabrer et al. 2017; Sullivan et al. 2018; Sedaghat et al. 2018), and, of course, light curve classification (Guiraud et al. 2019; Vithell et al. 2019) have benefited from the success of convolutional neural networks (CNNs) in image and sequential data processing. Autoencoders recurrent neural networks (RNNs), which preserve sequential information of input data, are found to be effective in extracting representative features from light curves (Waul et al. 2018). The encoder consists of the network takes light curves as inputs and generates features of the lower dimension. The decoder then attempts to reconstruct the light curves using only the encoder-generated features. By matching the reconstructed light curves with the original inputs, the autoencoder learns to isolate essential features that characterize the light curves. Unlike the conventional approach where the user and statistical features are hand-selected, the RNNs reduce feature extraction to an automatic process.

Most of the previous attempts at light curve classification focused only on correctly predicting object labels. A realistic application, however, is a distinguishable task is to correct objects with previously unseen types of variability, so called *novelty detection*, among a large number of objects of known classes. Zou et al. (2018) have recently proposed a deep network that combines autoencoder feature extraction with a Gaussian Mixture Model (GMM)-based supervised anomaly detection scheme.

In this letter, we present a neural network architecture that combines the efforts of Waul et al. (2018) and Zou et al. (2018). We jointly optimize an autoencoder RNN for feature extraction from variable light curves and an estimation network for classification and novelty detection. The motivation for this work is to promote the application of multi-task neural networks in variability analysis.

2. ALLS-SM ODS

2.1. Data and data pre-processing

The network is trained using the light curves from the ALLS-SM annotated Super for Superclass (ALLS-SM) Variable Stars Database and J1 (Saffee et al. 2014; Jayasinghe et al. 2018). The light curves, obtained

from the online database¹, typically have hundreds of observations and bands. From the test database, only light curves of types listed in Table 1 are selected for classification purposes. These variability types will be referred to as the variable *superclasses*.

The data selection and pre-processing procedure closely resembles that of Waul et al. (2018). To minimize confusion, we selected sources based on information from the ASAS-SN variable star catalog. Classification in the catalog, generated using the J1 classifier, are treated as the true labels of the sources. We note that these labels may not be completely genuine. Only variables with classification probabilities above 90% and at least 200 observations are used. Light curves with SUPERSMOOTHER residuals greater than 0.7 are omitted to ensure only true periodic sources are included. Following Jayasinghe et al. (2018), we further filtered out saturated and faint sources with $V < 11$ and $V > 17$.

In the ALLS-SM variable star database, sources that do not meet any of the classification criteria are referred to as 'variable stars of unspecified type' (VAR). This heterogeneous category contains objects with variability types that are not clearly differentiated in the superclasses. Since the light curves are by selection a large number of the known classes, the VAR objects are ideal for assessing the ability of the network in detecting unseen samples. Similarly, only VAR sources with at least 200 observational observations are selected. No classification probabilities or supervised residual cuts are applied to the light curve sequences in the VAR class.

Pre-processing began with the detrending of light curves to sequences of equal lengths, $n = 200$, as the autoencoder is optimized to process input sequences of fixed dimensions. The above pre-processing resulted in a reduced dataset of $\sim 8,000$ light curve sequences in the superclasses and $\sim 5,300$ in the VAR class. Using periods from the ALLS-SM catalog, the sequences are then phase-folded. The phase for each observational point, t , is then replaced by the relative phase between the current and the previous observation Δt . In particular, for the j -th measure, $\Delta t_j = t_j - t_{j-1}$; whereas $\Delta t_0 = 0$ is assumed for the first observation. The observed magnitudes in each light curve sequence x are normalized to have a mean of zero and a standard deviation of one, $x \rightarrow (x - \langle x \rangle)/s$, where $\langle x \rangle$ and s are the mean and standard deviation of the sequence. The measures are then normalized by the scale factor, $\sigma \rightarrow \sigma/s$.

2.2. Network Architecture

¹ <https://asas-sn.osu.edu/variables>

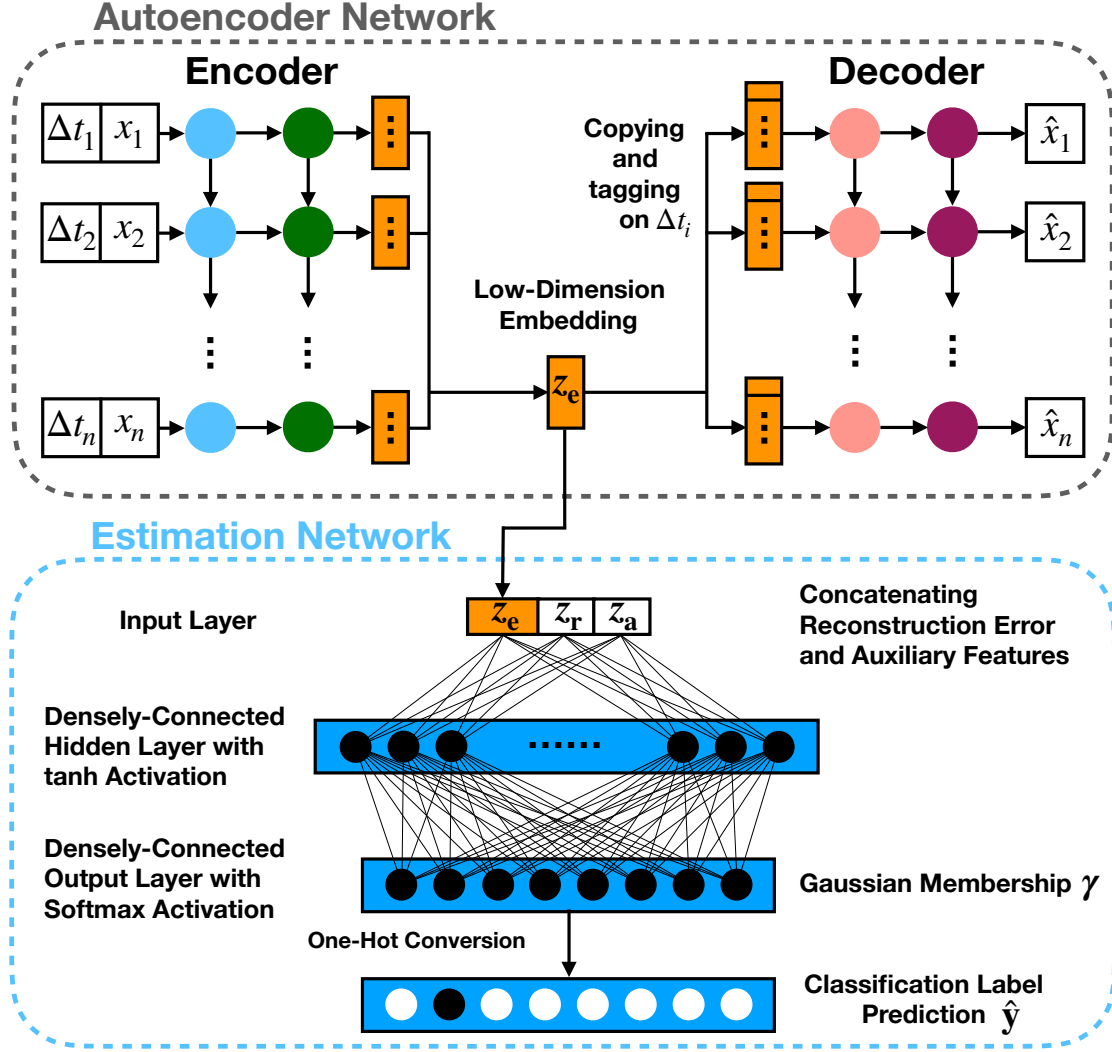


Figure 1. Schematic diagram of the neural network architecture. The structure of the DNN autoencoder follows Paul et al. (2018). The estimation network is similar to that used in Zeng et al. (2018) except for the added need of conversion for the classification task.

The schematic diagram of the neural network in this work is shown in Figure 1. It consists of two sub-networks:

Autoencoder network:

Feature extraction is performed using the autoencoder DNN in Paul et al. (2018). The encoder and decoder each consist of two layers of Gated Recurrent Units (GRUs).² A dropout layer² is included between the two recurrent layers to avoid overfitting (Srivastava et al. 2014). Given a batch of N normalized, phase-

folded input light curve sequences $(\Delta t_i, x_i, \sigma_i)$, where i denotes the i -th light curve, the encoder converts them to reduced-dimensional embedding vectors, $z_{e,i} \in \mathbb{R}^m$. The dimension of the embedding vectors, m , is a user-defined parameter called the embedding size. The embedding vectors are then used to reconstruct the input light curves $\hat{x}_i$ by the decoder. Following Paul et al. (2018), the loss function to be minimized is the weighted mean square error,

$$L_{AE} = \frac{1}{nN} \sum_{i=1}^N (x_i - \hat{x}_i)^2 \cdot w_i, \quad (1)$$

where $w_i = 1/\sigma_i$ is the sample weight for each observation, and the $(\cdot)^2$ is a common choice of activation. This loss function is advantageous over the conventional mean

² A dropout layer randomly ignores a fraction of units and their connections in a layer during training. It improves robustness by preventing the network from overfitting with sets of co-dependent weights. The reduction in network size also reduces training time.

Table 1. Variable superclasses and their constituent subclasses.

Variable Type	Superclass Abbreviation	S/S-Subclass ^a	Number of Light Curve Sequences ^b
Cepheids	CEP	CAB, CABB, DECP, DCEPS, DECPB	536
Delta Scuti Variables	DSCV	DSCV, DS	720
Eclipsing Binaries	EB	EB, EB, EB	25,713
Mira Variables	MI	MI	1,020
Potential Variables	POT	POT	805
PP Prae of Type A and B	PPAB	PPAB	7,809
PP Prae of Type C and D	PPCD	PPC, PPD	3,011
Semi-Regular Variables	SR	SR	8,156

^a Readers are referred to [Jagannathan et al. \(2018\)](#) and the S/S-Subclass Variable Stars online database for the detailed descriptions of the subclasses.

^b Number of light curve sequences after data pre-processing.

Quare error in that it explicitly includes measures to detect and reject irregularly sampled observations. Through minimization of the deviation between the input and reconstructed light curves, as defined by L_{AE} , the encoder is directed to reproduce bedded vectors that contain representative information of the original light curves.

Estimation network:

Classification and anomaly detection are simultaneously accomplished by the estimation network. The input feature vector to the estimation network is constructed by combining the autoencoder bedded vector $z_{e,i}$, two additional reconstructed error features $z_{r,i}$, and three auxiliary features $z_{a,i}$. Following [Zou et al. \(2018\)](#), the two reconstructed error features are the Euclidean distance and the cosine similarity,

$$z_{r,i} = \left[\frac{\|x_i - \hat{x}_i\|_2}{\|x_i\|_2}, \frac{x_i \cdot \hat{x}_i}{\|x_i\|_2 \|\hat{x}_i\|_2} \right], \quad (2)$$

where $\|\cdot\|_2$ denotes the L_2 norm. The mean and standard deviation of each light curve sequence are combined with the variable's period to form the auxiliary features, $z_{a,i} = (\langle x \rangle, s_i, \log_{10}(P_i))$. The input feature vector to the estimation network takes the form $z_i = (z_{e,i}, z_{r,i}, z_{a,i})$, with dimension, $l = m + 5$.

The estimation network connects the input feature vector to an output layer through a single dense connected hidden layer. The dimension of the output layer, K , is a pre-specified number of variable classes, which also corresponds to the number of Gaussian mixture components. A dropout layer is added after the hidden layer to minimize overfitting. The output layer adds a softmax activation function, producing a K -dimensional, normalized vector, γ_i , whose elements

can be interpreted as the probabilities of belonging to each variable class, Gaussian component.

The classification functionality is trained by minimizing the categorical cross-entropy loss,

$$L_{CE} = -\frac{1}{N} \sum_{i=1}^N y_i \cdot \log(\gamma_i), \quad (3)$$

where y_i is the true label of the i -th light curve expressed as a K -dimensional one-hot vector³, and the log is the natural logarithm or γ_i . To generate the classification label predictions, the softmax outputs, γ_i , are converted to one-hot vectors, $\hat{y}_i$.

The feature vectors, z_i , and the estimation network outputs, γ_i , are used to compute the Gaussian parameters using Equation (5) of [Zou et al. \(2018\)](#). Following the notation, μ_k is the mean location, Σ_k is the covariance matrix, and ϕ_k is the normalized Gaussian probability of the k -th Gaussian component in the feature space. The sample weight of each light curve sequence can be computed by

$$E(z_i) = -\log \left(\sum_{k=1}^K \phi_k \xi_k(z_i) \right), \quad (4)$$

where ξ_k is the normalized probability density in the k -th Gaussian component,

$$\xi_k(z_i) = \frac{\exp(-\frac{1}{2}(z_i - \mu_k)^T \Sigma_k^{-1}(z_i - \mu_k))}{\sqrt{(2\pi)^l |\Sigma_k|}}, \quad (5)$$

log is the natural logarithm, and $|\cdot|$ denotes the determinant. During network training, the Gaussian

³ A one-hot vector is an integer vector with all but one element set to zero, with the non-zero element having a value of unity at the location denoting its membership among one of the K classes.

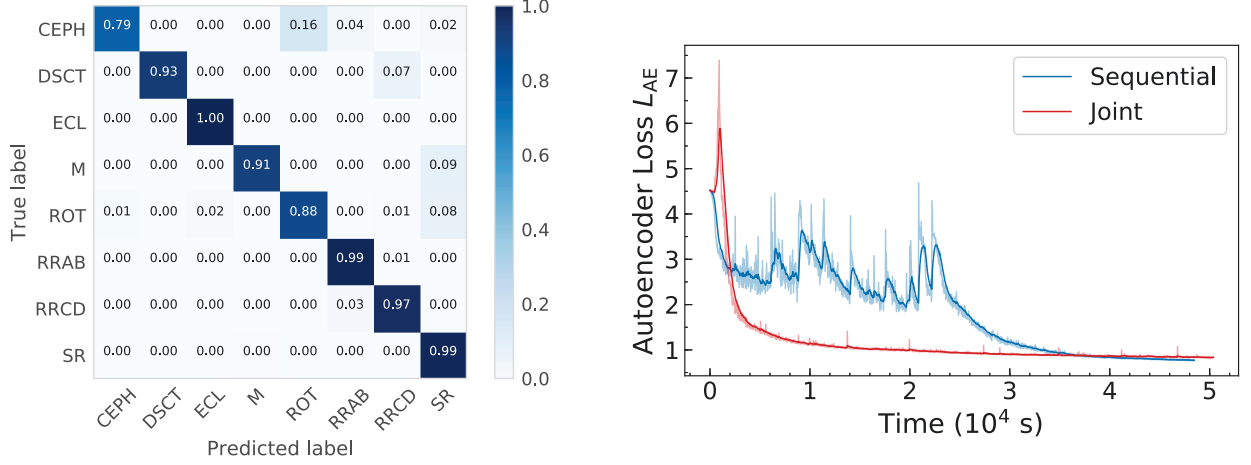


Figure 2. Left: Normalized confusion matrix for the jointly trained autoencoder-estimator. Right: The evolution of the autoencoder loss L_{AE} computed from the validation dataset during estimator training. The solid lines show the mean and the shaded areas show the standard deviation. The lines with lighter colors show the actual mean and standard deviation.

Parameters are determined and the associated loss is minimized using

$$L_{GMM} = \frac{1}{N} \sum_{i=1}^N E(z_i) \quad (6)$$

Optimizing the estimator network by minimizing L_{GMM} is equivalent to fitting the GMM parameters by maximizing the log-likelihood. After training, the GMM parameters should adequately describe the distribution of variable types in the feature space. We focus on variability can then be detected as outliers far away from the Gaussian components.

The overall estimator network architecture follows Zou et al. (2018) closely, but there are two main differences. First, the network is tasked with the additional problem of classification. Second, with the addition of the classification loss, the GMM is not susceptible to the singularity problem. The cross-entropy loss is therefore unnecessary and omitted.

2.3. Training Strategies

The dual network architecture was first introduced using the GasNet framework (Collet et al. 2015) for a solar outburst based (Adami et al. 2015). It was built on the GAN framework provided by Paul et al. (2018). The source code and a subset of the ASAS-SN data are available at <https://github.com/bhutsang/DeepClassifier novelty Detection>.

Weights and biases in both networks are initialized with the GLOROT_UNIFORM initializer (Glorot & Bengio 2010), and the loss function is minimized using the ADAM optimizer (Kingma & Ba 2014). We use the same eight variable subclasses as used in Nayashige et al. (2018, Figure 29) so that direct comparisons can be

made between classification performance. The variable subclasses in the above table are summarized in Table 1.

To demonstrate the versatility of the dual network, two training approaches have been carried out, namely, the joint and sequential training. In joint training, both the autoencoder and estimator networks are optimized simultaneously by minimizing the total loss,

$$L_{tot} = L_{AE} + (L_{GMM} + L_{CE}) \quad (7)$$

where the re-factor controls the relative importance of the GMM component. Default values of 1 are been tested and a default value of 10^{-3} works well for the current application. For sequential training, the autoencoder is first trained utilizing only the L_{AE} loss, i.e. the exact training approach used in Paul et al. (2018). The estimator network is trained afterwards, minimizing the parameterized loss terms of Equation (7). The motivation for including the sequentially trained models is twofold: it reproduces an already trained autoencoder with which we can better assess our classification accuracy, it also requires an opportunity to assess the possibility of attaching novelty detection functionality to re-trained feature extraction approaches.

A 80/20 split is used to divide the light curve sequences into training and validation datasets. For re-seeing the percentages of sequences in each variable class, the partition is performed using the STRATIFIEDKFOLD function of the Python SKLEARN package. The validation dataset is withheld from the estimator during the entire training stage, and is used to determine the classification accuracy of the estimator.

The following parameters are used in both training approaches. A total of 96 GPUs are used in both

Table 2. Classification accuracy and novelty detection performance for the reduced set of 16 hidden units per layer. Numbers in parentheses correspond to results from networks trained using the additional polynomial features.

		Joint Training	Sequential Training
Classification Accuracy	Estimator Network	98.8% (99.1%)	96.7% (96.6%)
	Random Forest	99.2% (99.4%)	99.2% (99.2%)
Novelty Detection Scores 95th Percentile Cutoff	Precision	0.885 (0.898)	0.875 (0.896)
	Recall	0.815 (0.957)	0.778 (0.933)
	F_1 score	0.85 (0.927)	0.826 (0.914)
Novelty Detection Scores 80th Percentile Cutoff	Precision	0.686 (0.700)	0.683 (0.696)
	Recall	0.91 (0.991)	0.935 (0.986)
	F_1 score	0.793 (0.821)	0.789 (0.816)

layers of the PPV autoencoder network. Training is done with a constant batch size of 2000 and learning rate of 2×10^{-4} for 2000 epochs. Dropout rates are fixed at 0.25 and 0.5 for the autoencoder and estimator networks, respectively. We varied the bedding sizes between 8, 16, 32, and 64. The size of the hidden layer in the estimator network is fixed at the bedding size. As a baseline, a separate grid of PPV classifiers are trained using the encoder-generated features after network training. We adopted a grid of n_estimators $\in \{50, 100, 250\}$, criterion $\in \{\text{gini, entropy}\}$, max_features $\in \{0.05, 0.1, 0.2, 0.3\}$, and min_samples_leaf $\in \{1, 2, 3\}$ in the SKLEARN `RANDOMFORESTCLASSIFIER` for evaluation.

3. RESULTS AND DISCUSSIONS

3.1. Classification Accuracy

As bedding sizes of 8, 32 and 64 only offer marginally different performance in both classification and novelty detection, in this paper, we will focus on the results from the reduced nodes with a bedding size of 16. The visualized confusion matrix from the joint network training is shown in the left panel of Figure 2. The overall classification accuracy of the validation data is 98.8%. In particular, the classification performance on classes with continuous outputs of light curves, namely EOP, PPV, and SP, is near perfect. In sequential training, the classification accuracy falls slightly to 96.7%, which is likely due to misclassifications in the less related superclasses.

As a comparison, the best-performing PPV classifier gives an accuracy of 99.2%, regardless of whether the network is trained jointly or sequentially. The high accuracy of PPV is expected given its complexity relative to the basic estimator network. In fact, the best-performing PPV classifiers in the grid typically contain

$\sim 10^2$ decision trees, each consists of $\sim 10^2$ nodes. A PPV classifier network contains $\sim 10^4$ trainable parameters, whereas the densely-connected layers of the estimator network have $\sim 200 - 5000$ depending on the bedding size. The classification accuracy and novelty detection performance scores for both training approaches are summarized in Table 2.

The right panel of Figure 2 shows the validation autoencoder loss over the duration of network training. The joint training approach offers a efficient training of the autoencoder, completing most of its learning early on after 10^4 s, or just ~ 200 epochs. Even though sequential training reduces slightly the autoencoder loss by the end of the training, classification accuracy is lower regardless. It suggests that by simultaneously optimizing both networks, the autoencoder is able to better retain class-discriminative information as an effective feature extractor. Across all bedding sizes (8, 16, 32, and 64), the joint training appeared to reduce stochastic fluctuations in the autoencoder loss, allowing superior consistent training.

3.2. Novelty Detection Performance

After the network training, the optimal parameters are computed and fixed using the entire training set. The VAR light curve sequences are fixed in the validation dataset to form the test dataset for novelty detection, which contains objects both from the superclasses and of the VAR type. Light curve sequences from the test dataset are then passed in to the encoder to generate the bedding vectors. Together with the reconstructed error and auxiliary features, the similarity of individual sequences can then be computed using Equation (4). Light curves in the similarity matrix above a pre-selected percentile cutoff in the training set, e.g. at 80% or 95%, are flagged as novel samples.

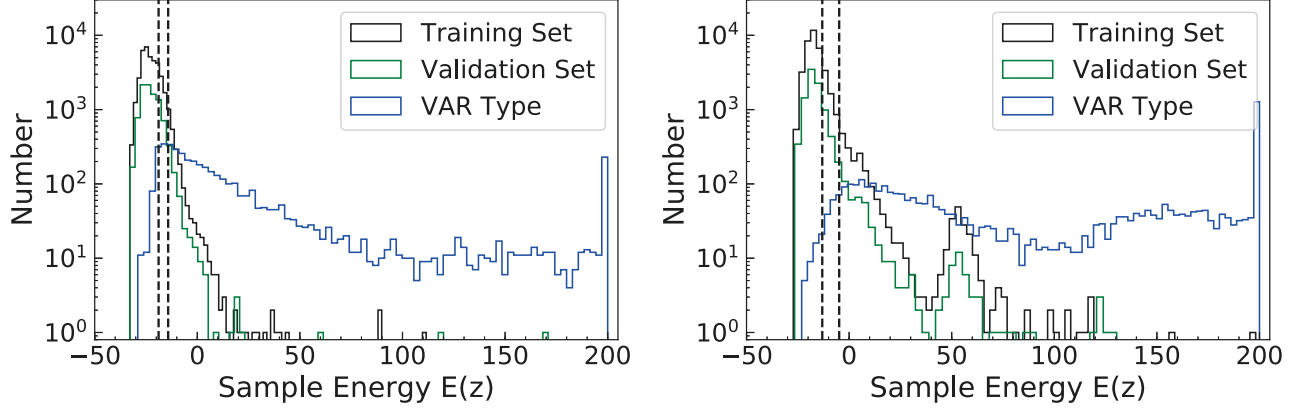


Figure 3. The energy histograms of the jointly trained set, split out (left) and split (right) the inclusion of photometric features. The energy cap of 200 is imposed for creating the histograms for better visualization. The vertical dashed lines denote the 80% and 95% percentiles of the training set. The secondary peaks at $E \sim 50$ in the case of the included photometric features are the result of missing magnitudes/colors, where a flag value of 99 is used.

The left panel of Figure 3 shows the energy histograms of the light curve sequences in the training/validation dataset and of the VAR type. Although there is a slight amount of overlap below $E = 0$, the majority of the VAR objects are higher in energy and are usually distinct. Using a 95th percentile cutoff, the precision, recall, and F_1 score are 0.885, 0.815, and 0.82, respectively. In the sequential training, the performance scores are slightly lower, at 0.875, 0.778, and 0.82, respectively. The reduction in recall relies on that the associated VAR samples yielded to samples with lower energies and not detected, suggesting that the G_{BP} colors are not as well-tuned. The performance scores are also summarized in Table 2. Despite the lower accuracy and novelty detection scores, the performance of sequential training is satisfactory. It suggests that the estimation of ϵ_{VAR} can be readily integrated with existing classification pipelines with pre-computed features, e.g. from Fourier analysis, to deliver additional novelty detection functionality.

3.3. Inclusion of Photometric Information

The majority of mis-classification in the confusion matrix (Figure 2) appears among classes CEP, V, and POX, whose light curves are scarce and thus statistically similar to one another. Though all classes benefit from explicitly including the auxiliary features (central magnitude, standard deviation, variable period), the three classes above require them for satisfactory classification accuracy. Since the estimation of ϵ_{VAR} is responsible for classification and novelty detection, it is expected that the required classification accuracy will allow G_{BP} colors to be better optimized for advanced novelty detection.

Photometric quantities from external catalogs are added as additional features in an attempt to reverse classification accuracies and thus require novelty detection. The photometric values are retrieved from the ASAS-SN Variable Stars Database and concatenated as part of the input features to the estimation of ϵ_{VAR} . As our current intention is to evaluate the feasibility of incorporating additional photometric features, only three catalog quantities are selected, namely, the G_{BP} magnitude, the $G_{BP} - G_{RP}$ color, and the G_{RP} band magnitude W_{RP} . The three confused classes are distinctly separated in this color-magnitude space (Lazarski et al. 2018, Figure 22). All parameters of the ϵ_{VAR} are extended to the training with the photometric features so a direct comparison can be made.

With the addition of the photometric information, the realized accuracy in the CEP, V, and POX classes are boosted up to 8%, 93%, 93% respectively. The overall classification accuracy increases up to 99.1%. Most importantly, the novelty detection performance is required to reach an F_1 score of 0.93. This exercise demonstrates that embedded light curve features can be trivially integrated with photometric features for better multi-task performance.

4. CONCLUSIONS

We present a dual-network architecture that allows simultaneous training for feature extraction, classification, and novel sample detection of variable star light curves. We have combined the recurrent neural network autoencoder for the series data reported by Waul et al. (2018) with the Gaussian mixture model detection network reported by Zeng et al. (2018). Applied to light curves from the ASAS-SN variable star

- Srinastha, V., Kishore, G., Krishnakumar, S., Sutskever, I., & Salathutdoo, P. 2018, *Journal of Active Learning Research*, 15, 1929.
<http://arxiv.org/papers/15/srinastha/180101>
- Zeng, B., Song, Q., & Li, P., et al. 2018, *International Conference on Learning Representations*.
<https://openreview.net/forum?id=Byjbb0->