

Challenges with Extreme Class-Imbalance and Temporal Coherence: A Study on Solar Flare Data

Azim Ahmadzadeh
dept. Computer Science
Georgia State University
Atlanta, GA, USA
aahmadzadeh1@cs.gsu.edu

Maxwell Hostetter
dept. Computer Science
Georgia State University
Atlanta, GA, USA
mhostetter1@cs.gsu.edu

Berkay Aydin
dept. Computer Science
Georgia State University
Atlanta, GA, USA
baydin2@cs.gsu.edu

Manolis K. Georgoulis
RCAAM of the Academy of Athens,
Athens, Greece
manolis.georgoulis@phy-astr.gsu.edu

Dustin J. Kempton
dept. Computer Science
Georgia State University
Atlanta, GA, USA
dkempton1@cs.gsu.edu

Sushant S. Mahajan
dept. Physics & Astronomy
Georgia State University
Atlanta, GA, USA
mahajan@astro.gsu.edu

Rafal Angryk
dept. Computer Science
Georgia State University
Atlanta, GA, USA
angryk@cs.gsu.edu

Abstract—In analyses of rare-events, regardless of the domain of application, class-imbalance issue is intrinsic. Although the challenges are known to data experts, their explicit impact on the analytic and the decisions made based on the findings are often overlooked. This is in particular prevalent in interdisciplinary research where the theoretical aspects are sometimes overshadowed by the challenges of the application. To show-case these undesirable impacts, we conduct a series of experiments on a recently created benchmark data, named Space Weather ANalytics for Solar Flares (SWAN-SF). This is a multivariate time series dataset of magnetic parameters of active regions. As a remedy for the imbalance issue, we study the impact of data manipulation (undersampling and oversampling) and model manipulation (using class weights). Furthermore, we bring to focus the auto-correlation of time series that is inherited from the use of sliding window for monitoring flares' history. Temporal coherence, as we call this phenomenon, invalidates the randomness assumption, thus impacting all sampling practices including different cross-validation techniques. We illustrate how failing to notice this concept could give an artificial boost in the forecast performance and result in misleading findings. Throughout this study we utilized Support Vector Machine as a classifier, and True Skill Statistics as a verification metric for comparison of experiments. We conclude our work by specifying the correct practice in each case, and we hope that this study could benefit researchers in other domains where time series of rare events are of interest.

Index Terms—class imbalance, sampling, time series, flare forecast

I. INTRODUCTION

To gain valuable insights or robust predictive performance from data, we must first ensure the integrity of our data. Beyond data collection, this involves data-cleaning. It requires a thorough investigation by the experts of the domain and data scientists to produce a reliable dataset. Nonetheless, there are some challenges which are inherited from the subject under study due to unique characteristics of the data which should be identified, understood and dealt with appropriately. Class-imbalance issue is one of the main problems of this kind,

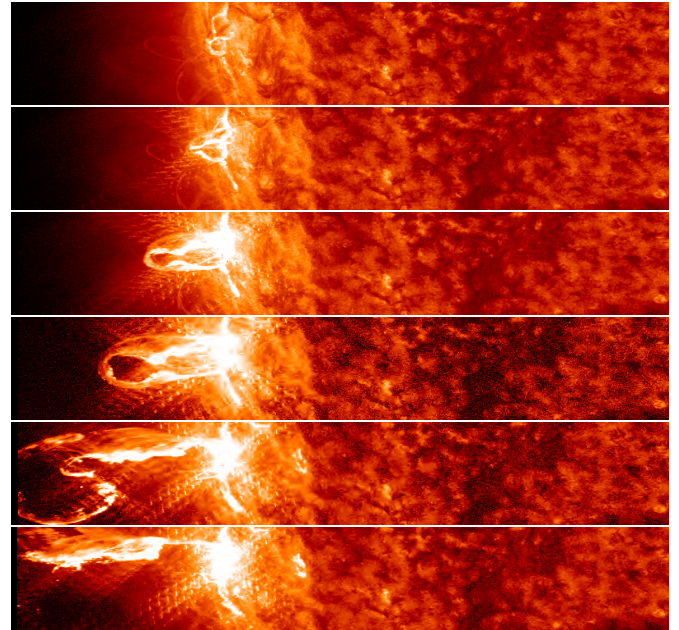


Fig. 1. Snapshots of an X-class flare, peaking at 7:49 p.m. EST on Feb. 24, 2014, observed by NASA's Solar Dynamics Observatory, in the 304Å wavelength channel. (Images source: <https://helioviewer.org/>)

which is present in many natural or other nonlinear dynamical systems. This is often due to the nature of the events, not the data collection process.

Class-imbalance is a common problem, with many potential remedies. Some of these remedies are common and well-known, but can still be misapplied. This is particularly true when the primary objective is not machine learning per se but the testing and scrutiny of domain-specific theories. The complexity of the problem at hand and the absence of data experts very often underestimate the needed level of care,

resulting in unrealistic and unreliable analyses, with little practical value.

In this study, we use a dataset of extracted statistical features from the time-series of solar magnetic data as an example for extreme class-imbalance. We show the impact of disregarding an interesting phenomenon called “temporal coherence” on classification of events in spatiotemporal datasets. We demonstrate the impact and biases of different class-imbalance remedies and discuss how they should be interpreted from the perspective of the subject under study by involving domain experts. We hope that this work raises awareness to interdisciplinary researchers and enables them to spot and tackle similar problems in their respective areas.

All our data manipulation steps, such as preprocessing, feature extraction, dimensionality reduction, data normalization, and sampling, are made publicly available as a general-purpose python module [1] that is still an active project.

II. COMMON CHALLENGES IN SOLAR FLARE PREDICTION

Machine learning applications to solar flare forecasting are not new, but we have yet to see any significant forecast improvement resulting from these efforts. In the following sections, we present a few challenges commonly faced while tackling solar flare prediction problem from a data-driven perspective. Predicting or forecasting¹ the occurrence of solar flares is a typical 21st century rare-event task. Flares are sudden and substantial enhancements of high energy electromagnetic radiation (like extreme ultra violet and X-rays) at local solar scales which pose a threat to humans and equipment in space. They are automatically detected and classified by the National Oceanic and Atmospheric Administration’s (NOAA) constellation of GOES satellites based on their peak flux in soft X-ray wavelengths on a logarithmic scale as A-, B-, C-, M- and X-class solar flares. A and B-class flares are difficult to distinguish from the random variations in the Sun’s background X-ray level, but C-class flares and above are detected reliably by GOES satellites. The most intense of these classes, namely M and X, are most often targeted for prediction due to their potentially adverse space-weather ramifications. An example of an X-class flare is shown in Fig. 1.

A. Extreme Class-Imbalance

The frequency distribution of the peak X-ray fluxes of flares is nearly a perfect power law with a dynamical range spanning several orders of magnitude. A statistical analysis of NOAA’s flare reports during solar cycle 23 (1995 to 2008) shows that around 50% of active regions produce C-class flares, while 10% produce M-class flares and less than 2% produce X-class flares. Solar cycle 24 (2009 to present), from

¹In this document, we use the terms ‘forecast’ and ‘prediction’ interchangeably, as the data of this study are time series where ‘forecast’ seems to be a better fit, while our employed statistical approach towards forecast of flares is simply a ‘classification’ task, which is done on the history of flares. Since the classification determines the presence of flares in a fixed prediction window in future, the term ‘prediction’ is also appropriate. This decision also allows us to avoid repetition of the word ‘forecast’.

which our data discussed in Section 3 are taken [2], exhibit a much weaker major flare crop, making class-imbalance a conspicuous problem to deal with (for a review, see [3]).

B. Point-in-time vs Time Series Forecasting

Solar flares are significant precursors to other space weather phenomena. Much of the existing literature (e.g., [4], [5]) is focused on using instantaneous measurements to produce a binary or probabilistic flare forecast over a preset forecast horizon. However, flares are inherently dynamical phenomena, with clear pre-flare and post-flare phases, characterized by certain evolutionary trends [6], [7]. Because of this, time series of aspiring flare forecasting parameters should be used, rather than isolated points in time. We believe that the sheer level of difficulty of this task was the key factor for the (over-)simplifying point-in-time assumption. However, it may be precisely this compromise that may have hampered non-incremental progress toward flare prediction. Therefore, our main goal in this study is to explore the difficulties which arise when we take into account the temporal evolution of magnetic parameters of active regions rather than looking at a single snapshot in time.

C. Non-representative Datasets

We have access to data from many scientific instruments that map the magnetic field of the sun’s photosphere. This has been monitored for 25 years [8]–[12], but the time span for high-quality data is still limited. Another issue for machine learning applications is that training forecasts on one part of the solar cycle is not necessarily effective for forecasting other parts of the solar cycle. This is because of continuously modulating background of magnetic activity. Some of these effects are evident in the experiments in Sec. V. Therefore, there exists a problem of heterogeneous and/or non-representative data. Coupled with simpler, yet unjustified, selections of random undersampling from the majority class events, the sampled subsets of data fail to become representative of the overall flare population.

III. SWAN-SF DATASET: A MULTIVARIATE TIME SERIES DATA

We are only just recently getting access to high-quality time-series data for use in solar flare prediction. Many previous and current projects use point-in-time measurements [4], [5], [13]–[15]. It is possible that a time-series approach will enable new progress on using machine learning for flare prediction. Here we will use a benchmark dataset, named as Space Weather ANalytics for Solar Flares (SWAN-SF), recently released by [2], and made entirely of multivariate time series, aiming to carry out an unbiased flare forecasting and hopefully set the above question to rest.

The SWAN-SF dataset is made of five partitions. These partitions are temporally non-overlapping and divided in such a way that each of them contains approximately an equal number of X- and M- class flares (see Fig. 2 and 3). The data points in this dataset are time series slices of physical

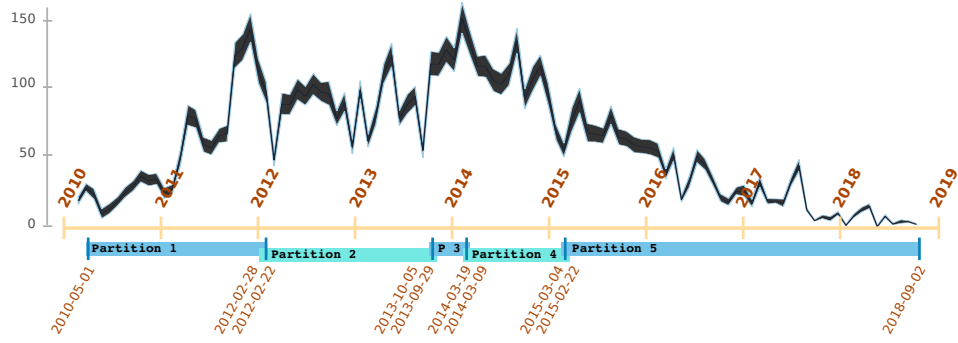


Fig. 2. Time span considered for each partition of SWAN-SF dataset, the monthly report of average number of sunspots per day, and the daily variance represented with a gray ribbon (Source: WDC-SILSO, Royal Observatory of Belgium, Brussels).

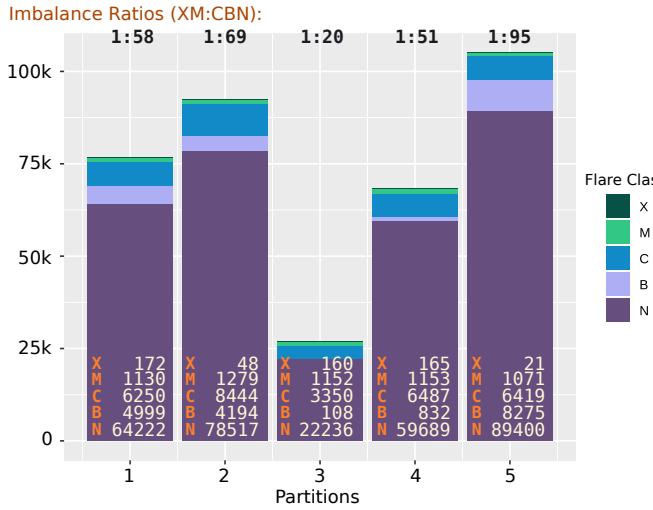


Fig. 3. Counts of the five flare classes across different partitions.

(magnetic field) parameters extracted from the flaring and flare-quiet regions, in a sliding fashion. That is, for a particular flare with a unique id, k equal-length multivariate time series are collected from a fixed period of time in the history of that flare. This period is called an *observation window*, denoted by T_{obs} , and spans over 24 hours. Given that t_i indicates the starting point of the i -th slice of the multivariate time series, the $(i + 1)$ -th slice starts at $t_i + \tau$, where $T_{obs} = 8\tau$. This kind of sliding observation inherits the fact that very often the k slices behave very similar due to their temporal and spatial closeness. In other words, the physical parameters describing the behavior of the region corresponding to a particular flare are not significantly different from one slice to the next. These similar slices, if described in the multi-dimensional feature space of our data, are located too close to each other to be considered distinct instances. And their closeness does not reflect any characteristics of those data points, except our slicing methodology. We refer to this phenomenon as *temporal coherence*² of data. This is a key concept to understand some

²We introduce this concept in the context of data manipulation, and it should not be confused with *temporal coherence* in Optics or any other topic.

of the challenges we would like to address in this study.

A. Final Forecast Dataset

When working with time-series, we can take one of two main approaches. We can preprocess and use the data directly, with methods optimized for time-series. We can also extract a set of statistical features from the time-series and then use the extracted descriptors in our models. Our interest in the analysis of different sampling methods guides us to use the extracted features.

Statistical Features. Though, as mentioned in Sec. II-B, point-in-time measurements have been proposed for flare prediction ([4], [5]), there are no established theories or works on which time-series characteristics may be useful for flare forecasting. We hereby build a prediction dataset relying on the set of first four statistical moments of the time series, namely, their *mean*, *variance*, *skewness*, and *kurtosis*. To allow comparison with previous studies, we also consider a point-in-time feature, namely the *last value* of each time series (from the observation window). Moreover, we also add *median* to the list to compensate for the effect of outliers on *mean*.

The obtained dataset of the extracted features has a dimensionality of 144, resulting from the computation of the 6 above-mentioned statistics on the 24 physical parameters of the SWAN-SF dataset. Data points of this dataset are labeled by 5 different classes of flares, namely GOES X, M, C, B, and N. The latter represents flare-quiet instances or GOES A-class events.

This work is focused mostly on concerns with techniques common to machine learning applications, to help describe a robust methodology independent of data or model specifics. To that end, we use the *last value* which is most comparable to other point-in-time studies. At the end, however, we present the contribution of other statistical features as well to show the benefit of using time series instead of point-in-time data instances. We also conduct our experiments on a binary class data by merging X and M classes into a super-class called XM, and C, B, and N classes into another super-class denoted by CBN. This simplification allows us to only focus on the challenges we mentioned before, which is the primary objective of this study.

Preprocessing. After computing the above-mentioned features, the dataset requires a minimal preprocessing due to the presence of some missing values. Since this accounts for a very small fraction of the data (i.e., $< 0.01\%$), we simply utilize linear interpolation to fill out those values. In addition, we use zero-one data transformation to normalize the data for our experiments, since otherwise the optimal hyperplanes found by SVM would be meaningless.

IV. CLASS-IMBALANCE AND TEMPORAL COHERENCE

In this section, we discuss two problems that are common in machine learning and always present in data used for solar flare prediction: class-imbalance and temporal coherence. Without loss of generality, we use the SVM classifier which, like many other learners, is sensitive to these issues. Therefore, instead of analyzing the specific characteristics of each classifier, we focus on the common denominators of the well-known family of classifiers in view of imbalanced datasets.

A. Class-imbalance Problem

Class-imbalance describes a case where the occurrences of one or more of the class(es) in the data are far less than that of the other class(es). When class imbalance is present, it is well-established to apply special techniques to deal with it, since generally machine learning models perform better when classes are roughly equally frequent. Here we use the terms “minority”, or “positive”, class to refer to the less frequent group and “majority”, or “negative”, class, conversely. Stronger flares (GOES X- and M- class) form our minority class, and weaker events (GOES C-, B- and A-class) belong to our majority class. From the practical point of view, we care significantly more about the minority class, since stronger flares have much bigger impact on our environment. This is why we refer to this class as ‘positive’. Fig. 3 illustrates the distribution of all classes in each partition.

In classification, the aim is to optimize our objective function by minimizing the number of misclassifications. When severe class-imbalance is present, often correct classifications of the minority class are more valuable than correct classifications of the majority class. Since the density of the majority class is significantly higher than that of the minority class, many instances of the majority class may be sacrificed (i.e., misclassified) for a correct classification of an instance from the minority class. The SVM classifier in particular, searches for optimal hyper-planes to make such separations. An imbalanced dataset most likely preserves the imbalanced density of the classes even close to the decision boundaries (where the ideal class regions overlap or meet). In such a situation, a hyper-plane that is supposed to pass through the boundaries will be shifted into the region of the minority class to reduce the total number of incorrect classifications/predictions by getting all the positive classes right. This leads to higher true-negatives (i.e., correct predictions of CBN-class flares) and lower true-positives (i.e., correct predictions of XM-class flares). In other words, a model in a class-imbalance space always favors the majority class. This is of particular concern because virtually

all real-life, class-imbalance problems aim to predict minority, rather than majority, events.

The effects of class-imbalance are also seen during forecast performance evaluation. Many well-known performance metrics are significantly impacted by class-imbalance, including accuracy, precision (but not recall), and the f1-score. This is mainly because these measures ignore the number of misclassifications. For instance, on imbalance datasets, a model that classifies all instances as the negative (majority) class may result in a very high (often asymptotic to 1.0) accuracy, while learning little or nothing about the minority class. For the particular case of class-imbalance there are defined less susceptible measures such as TSS (True Skill Statistic³ [16]) or the HSS (Heidke Skill Score [17], [18]), with TSS being reportedly more robust for solar flare prediction [14]. In this study, we use only TSS since our main objective is to show the changes in the models’ performance and not to find an operation-ready model.

Undersampling and Oversampling. A well-known approach to handling class-imbalance is to make our training dataset balanced. This is often done by *undersampling* (that is, taking out instances from the majority class) or *oversampling* (that is, providing more instances to the minority class by replication). Undersampling results in using only as many negative instances (majority) as there are positive instances (minority) in the training phase, thus achieving a 1:1 balance ratio. This solution, however, comes at some cost. When undersampling, we leave out a great portion of the data during training, therefore not learning from the entire collection. To avoid the enormous data waste, a very large dataset should be available overall. When oversampling, we add replicates of the existing instances. This may cause a model to memorize the patterns and structures of minority events in the data instead of generalizing and learning about them; a bad practice that is very prone to overfitting.

These sampling methodologies are often used and can be effective for improving classification performance, but their applications should be specific to our data and align with the objectives of our work. One should be extra careful when applying them to a multi-class data, such as the flare dataset. This remains true despite the fact that we converted it to a binary class problem. In Table I, we list a set of sampling methodologies, which are also visualized in Fig. 4. As it is shown, alternative avenues exist depending on whether a 1:1 balance in the sub-class level (i.e., $|X| = |M|$ and $|C| = |B| = |N|$) is also required or not. Notice that this is an addition to the primary goal of our sampling which aims to achieve a balance between the super-classes (i.e., $|XM| = |CBN|$). When undersampling, for instance, if this additional balance is expected, we first need to decide which class in the minority group is considered the “base” class. Letting GOES X be the base class (see US2 in the table and the visualization), we must undersample from M-class flares first to balance X and M classes and then undersample from

³This is also known as Hanssen-Kuipers Discriminant.

TABLE I
DIFFERENT UNDERSAMPLING AND OVERSAMPLING APPROACHES.

Method	Abbr.	Description
Undersampling 1	US1	undersamples C , B , and N classes while preserving their original proportions.
Undersampling 2	US2	undersamples C , B , and N classes to achieve a 1:1 balance in sub-class level, as well as the super-class level, while keeping X class unchanged.
Undersampling 3	US3	undersamples C , B , and N classes to achieve a 1:1 balance in sub-class level, while keeping M class unchanged. This requires oversampling of X class.
Oversampling 1	OS1	oversamples X and M classes while preserving their ratios.
Oversampling 2	OS2	similar to OS1, except that this suppresses N -class by shrinking it down to $3 C - (C + B)$.
Oversampling 3	OS3	oversamples X and M classes to achieve a 1:1 balance in sub-class level, while keeping C class unchanged. This requires undersampling of B and N classes.
Oversampling 4	OS4	oversamples all five classes to achieve a 1:1 balance in sub-class level, while keeping N class unchanged.

the majority (CBN) class. This yields a balanced dataset in both super-class and sub-class levels. A quick look at other approaches shows that the choice of the sampling method plays a critical role. Knowing that GOES C class represents the strongest flares in the weak-flare class (majority), it is expected that a higher fraction of C-class instances in this group results in a harder prediction problem for the model. In other words, a sampling methodology has to contort the climatology of flares to achieve the desired balance. This change affects the distribution of samples in the feature space by making the decision boundary (where GOES C and M classes overlap) denser or sparser. As an example, comparing US2 or US3 with US1 where the C-class accounts for the smallest fraction, indicates that US1 yields a slightly simpler classification problem. Having that said, US1 would still be a realistic sampling since the climatology of flares is preserved within the majority class.

The above suggests that performance of different models on the same dataset is only comparable if they all employ identical sampling methodologies.

Misclassification Weight. We can also modify our objective function to penalize misclassifications of the minority class more than the majority. SVM, like some other machine learning algorithms, can incorporate different weights in its objective function. For details on how this is mathematically implemented, we refer the interested readers to [19]. For the experiments in this study, inspired by the class-imbalance ratio, we adjust the weights using $w_j = \frac{n}{k \cdot n_j}$, where n is the total population, n_j is the population in class j , and k represents the number of classes.

B. Cross Validation

Cross validation is a family of statistical techniques, typically used to determine the generalization power of a model. Regardless of the employed cross validation technique (k -fold, leave- p -out, stratified or purely random), it is very often assumed that a random selection is allowed. However, in many

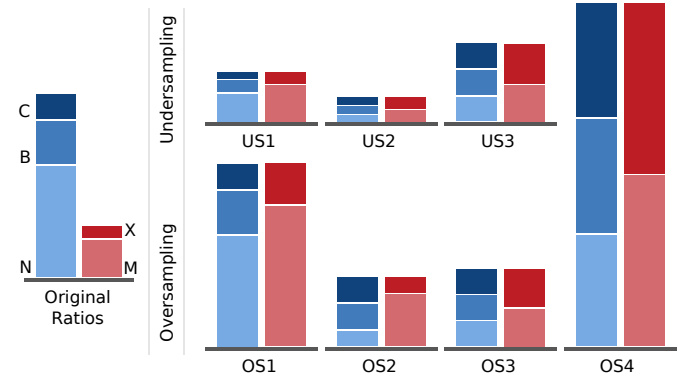


Fig. 4. Visualization of different undersampling and oversampling approaches. On the left, the original proportion of the sub-classes (X , M , C , B , and N) and super-classes (XM and CBN) is depicted. Each plot on the right corresponds to one sampling approach, i.e., a single row in Table I. This is a generic illustration and does not reflect the exact imbalance ratios in SWAN-SF.

real-world data collections, random sampling must take into account the spatiotemporal characteristics of the data, which is rooted in the temporal (or spatial) coherence of data. The temporal coherence in SWAN-SF dataset, as we discussed in Sec. III, prohibits this practice, since it yields an overly optimistic performance of the model’s learning; a phenomenon known as ‘overfitting’.

To avoid this mistake, we select training and testing instances from different partitions of dataset. To showcase the impact of randomly splitting the data from one partition, we conduct an experiment that is presented in Sec. V-C.

Validation Set. Disregarding the temporal coherence of the data and random sub-sampling for obtaining the training and testing sets has an additional and perhaps a more important impact: when the test set is obtained by randomly splitting the data, the tested model obscures the overfitting sign, i.e., a significant difference between the training and testing performance. Therefore, the use of any sorts of sampling methodology on the test set (either for reducing the imbalance ratio, as we discussed before, or in cross validation) must be avoided at all costs if it distorts the actual distribution of the events. Cross validation is also used for optimization of models’ hyper-parameters or the data-driven parameters. We wish to reiterate that for all such tasks, another independent subset of data should be used which is known as the validation set, and the test set must never be exposed to the model except for reporting the final performance of the model. For instance, to tune SVM’s hyper-parameters, namely c and γ for achieving an optimal hyper-plane by the model, a validation set must be defined in order to reflect the changes that the model takes in. Based on our expectation of the acceptable performance, we may let the model improve upon the validation set’s feedback. Only when we believe that the model has reached its highest performance we can use the test set to measure its robustness.

Similar to the capital impact of different sampling methodologies on performance, as we discussed previously, sampling modification of the test set is another way of creating non-

robust models with seemingly high performance.

C. Data Normalization

It is common, and depending upon our model, sometimes necessary to transform the ranges of all feature values into a unified range before training. This is called normalization and is done for a variety of reasons, most importantly for allowing the model to distinguish between different structures without taking into account the different units of measurements used for each feature. Simplicity of this concept makes researchers to sometimes overlook the variable impact of different ways of normalization. Regardless of which transformation function one may use, normalization can be done either locally and globally. Although it is a common practice to only take into account the global statistics (e.g., $\min$ and $\max$ of each time series in the entire dataset), for the cases similar to SWAN-SF dataset, this must be put to the test. Looking at our dataset, among the five partitions spanning over a period of ten years, the changes in solar magnetic activity from minimum to maximum during the solar cycle makes the statistical characteristics of different partitions significantly different. In such cases, a local normalization would impact some features negatively, or even generate some out-of-range values. We investigate this via another experiment in Sec. V-D, by showing the significantly different performance levels using different statistics.

D. Hyper-parameter Tuning

When performing hyper-parameter tuning, we should take temporal coherence into account. Any supervised learning model requires optimization of its hyper-parameters in a data-driven manner. Since the discovered hyper-parameters should remain optimal over the entire dataset (including the upcoming data points for a deployed model) the training and validation sets, as well as the test set, must each be representative of the entire dataset. In our dataset, because of the temporal coherence, random sampling does not produce such subsets. Hence, tuning process remains confined to the partitions. Although the solution we proposed for cross validation using NOAA AR Number, may also be used to tackle this problem, it is very likely that a model highly optimized on one partition simply may be not performing well globally. Therefore, we believe that this is a problem yet to be investigated more thoroughly since at this point it is very clear to us that flare forecast problem has a dynamic and periodic behavior, for which ensemble models may be more appropriate.

V. EXPERIMENTS, RESULTS, AND INTERPRETATIONS

In this section, we present a series of experiments, each designed to address one of the problems we have previously discussed. We also elaborate on the interpretation of these experiments in regards to our overarching flare forecasting task.

Our objective is not to train the next best flare forecasting model, but to compare models that were trained differently, to help establish data driven practices recommended for solar

flare forecasting. Therefore, although any changes on the data (e.g., using different normalization, data split, or sampling techniques) require re-tuning of the hyper-parameters, without loss of generality we rely on our pre-tuned hyper-parameters for SVM: $c = 1000$, $\gamma = 0.01$, with a Radial Based Function (RBF) kernel.

We use TSS to evaluate the results of each experiment only to show the changes different treatments cause. A high TSS does not necessarily indicate a robust forecast and is often only presented with an accompanying HSS. We do not present the HSS because we evaluate our results on all different partitions of our dataset which do not have equal levels of class-imbalance. HSS is susceptible to changes in class-imbalance and is not ideal for comparison between trials in our case.

A. Baseline

We establish a baseline performance by training the model without applying any class-imbalance remedies.

Experiment Z: Baseline. This experiment is as simple as training SVM on all instances of one partition and testing the model on another partition. We try this on all possible partition pairs, resulting in 20 different trials, to illustrate how the difficulty of the prediction task varies as the partitions are chosen from different parts of the solar cycle. The results are visualized in Fig. 5, along with the impact of discussed class-imbalance remedies that we further discuss in the following sections.

B. Tackling Class-imbalance Issue

In Sec. IV-A, we discussed three different approaches towards tackling the class-imbalance problem. To show the impact of each solution, we carry out some experiments and discuss the results below.

In the following experiments, we train and test SVM on all 20 permutations of the partition pairs independently. In each round, the model learns from instances in the training partition and is then tested against the (different) testing partition. To measure the confidence of a model's performance when a certain sampling method is employed, we repeat the experiment 10 times and report the variance and mean value of TSS.

Experiment A: Undersampling. During the training phase, we apply Undersampling 2 (US2 from Table I) which is an X-class based undersampling. This enforces a 1 : 1 balance, not only in the super-class level (i.e., $|XM| = |CBN|$) but also in the sub-class level (i.e., $|X| = |M|$ and $|C| = |B| = |N|$). The trained model is then tested against all other partitions one by one to examine the robustness of the model. The undersampling step is only taken in the training partition, as undersampling of the test partition distorts reality and would not reflect the true model's performance. The consistent and significant impact of this remedy is evident in Fig. 5.

Experiment B: Oversampling. Here, we use Oversampling 3 (OS3 of Table I) and perform an experiment similar to experiment A. Again, no over- or undersampling takes place

in the testing set. Comparing the results of oversampling with Undersampling 2, in Fig. 5, shows a close correspondence between the two models in terms of their mean TSS values; typically, differences are within applicable uncertainties.

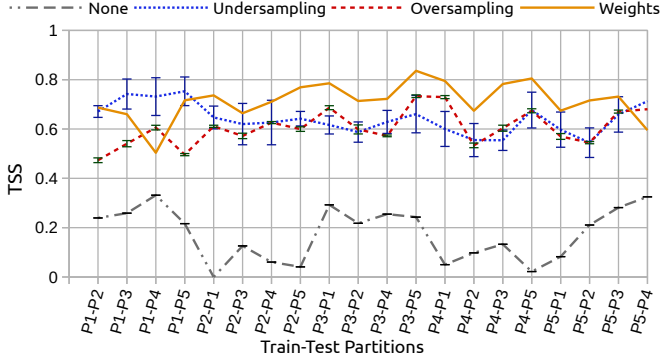


Fig. 5. [Experiments Z, A, B, and C] Average TSS of SVM trained and tested on all possible permutations of partition pairs, using three different remedies for the class-imbalance issue (undersampling, oversampling, and misclassification weights), compared to a baseline performance where no remedy is employed. In all cases, global normalization is used.

Experiment C: Mis-classification Weights. Here, we apply weights as discussed in Sec. IV-A. We use the imbalance ratio of the super-classes as the weights. For instance when working with *Partition 3*, since the minority-to-majority ratio is 1 : 20, we set $w_{XM} = 20$ and $w_{CBN} = 1$. As shown in Fig. 5, this solution outperforms both undersampling and oversampling approaches in terms of their TSS. It is worth pointing out that employing misclassification weights has the advantage of a data-driven tunability that may be better suited than over- and undersampling to achieve more robust forecast models.

C. Impact of Cross Validation

In this experiment we demonstrate the importance of temporal coherence and the potential pitfalls of random sampling that we discussed in Sec. IV-B.

Experiment D: Data Splits. To demonstrate the overfitting that occurs when we do not account for the temporal coherence, we train and test SVM on instances randomly chosen from the same partition. Technically, this is a k -fold cross validation using a random sub-sampling method with $k = 10$. The results are then juxtaposed with those obtained by training SVM on one partition and testing it on another. We equipped SVM in both scenarios with misclassification weights (the same as in Experiment C), to eliminate the need for an additional sampling layer. Therefore, the only determining factor is whether the instances are sampled from the same partition or not. Let it be clear that sampling from a single partition does not mean any overlapping between the training and testing sets.

The results are summarized in Fig. 6. When SVM is trained and tested on a single partition (denoted as ‘Unifold’ in the figure), performance is boosted very significantly with $TSS \in (0.84, 0.94)$ with an average $TSS \approx 0.88$. Training and testing on different partitions (denoted as ‘Multifold’ in the figure)

yields $TSS \in (0.50, 0.84)$ with an average $TSS \approx 0.71$. This remarkable difference should not be viewed as evidence of the robustness of the model but rather points to memorization and hence overfitting, caused when a forecast model is trained and tested on a temporally coherent dataset. It is the lower performance when the model is trained and tested on different partitions that better encapsulates its true robustness.

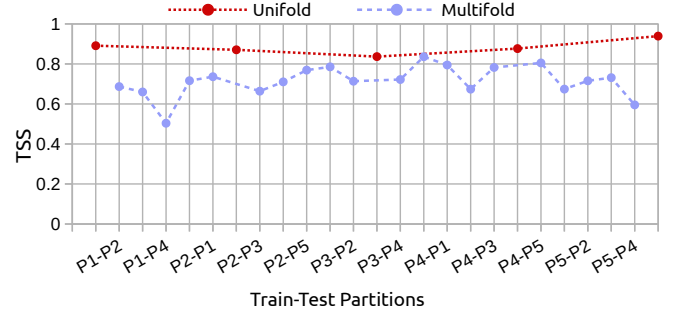


Fig. 6. [Experiment D] Average TSS of SVM performance when trained and tested either using a 10-fold cross validation sub-sampling on a single partition (unifold) or assigning different partitions for training and testing (multifold). In all cases, global normalization is used and the SVM has been equipped with misclassification weights.

D. Impact of Normalization

To demonstrate the impact of the different normalization techniques discussed in Sec. IV-C, we conducted the following experiment.

Experiment E: Normalization. To transform the feature space to a normalized space (using a zero-one normalization method) we apply global normalization on a pair of partitions by using global extrema ($\min$ and $\max$). We then train SVM on one partition and test on the other. In another attempt, we apply normalization on each partition separately, using the local extrema of the corresponding partition. In order to minimize the impact of other factors on our experiment, we avoid employing any remedies for handling the class-imbalance issue.

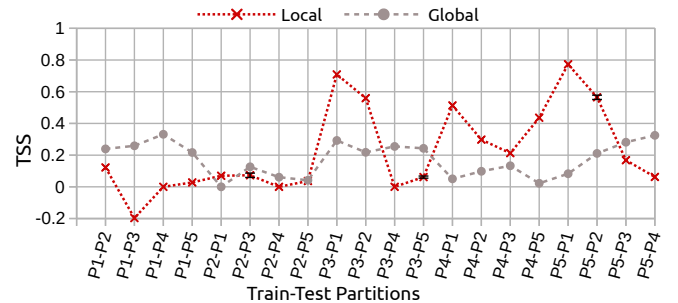


Fig. 7. [Experiment E] Average TSS of SVM performance impacted by two different data normalization approaches; global and local. In all 20 trials shown, no special remedy to the class-imbalance problem has been employed.

In Fig. 7, we can see dramatic differences between these two applications of normalization. This is because of the

continuously modulating magnetic activity during the 11-year solar cycle. This makes the normalization inconsistent between the training and the testing partitions in case of local normalization. Careful observation reveals that, when local normalization is utilized, it seems training on the second half of the solar cycle (*Partitions* 3, 4, and 5) builds models which perform better when tested on the instances from the first half of the cycle (*Partitions* 1 and 2). Conversely, when models are trained on the first two partitions and tested on the rest, their TSS drops significantly. Both of these observations strongly agree on the noticeable differences in data collected from different windows of time. On the other hand, global normalization leads to more robust models, as the changes in TSS values are relatively lower, compared to those in local normalization. This is achieved at the cost of a sizable decrease in the overall performance. While additional investigations of this effect are clearly warranted, one gathers that performance may be optimized when training and testing take place on similar levels of solar activity, when local and global normalization converge on similar normalized values. This, however, is hardly tenable in operational settings, when forecast models are standing by to perform on yet unknown day-by-day solar activity levels.

E. Oversampling Impact

In Sec. IV-A, we showed that there are multiple variants of oversampling and undersampling approaches. We also presented how this affects flare distributions in *Partition* 3 as an example. Below we test how different oversampling techniques impact TSS values across different partitions.

Experiment F: Oversampling with or without Sub-Class Balance. We use Oversampling 1 and 3 (from Table I) in the training phase to remedy the class-imbalance problem and then we test the trained model against all other partitions. We chose OS1 and OS3 since their differences make them an interesting pair; OS1, compared to OS3, replicates M-class instances with a significantly larger factor than it does with X-class instances, allowing a large number of ‘easier’ instances in terms of classification. Therefore, it is naturally expected that OS1 results in an ‘easier’ dataset, hence a higher classification performance. Here we use global normalization without meaning to imply that it performs necessarily better. Our results are shown in Fig. 8. For them, one sees a relatively similar, consistent performance, although the climatology-preserving Oversampling 1 gives a statistically higher performance. This said, it becomes clear that different oversampling methods give non-identical performances. Therefore, comparison of any two forecasting models on similar datasets will be fair only if the employed sampling methodologies are identical.

F. Using Other Time Series Features

As mentioned in Sec. III-A, all previous experiments involved using only the *last value* feature. In this experiment, we present one of the benefits of using time-series as opposed to point-in-time measurements.

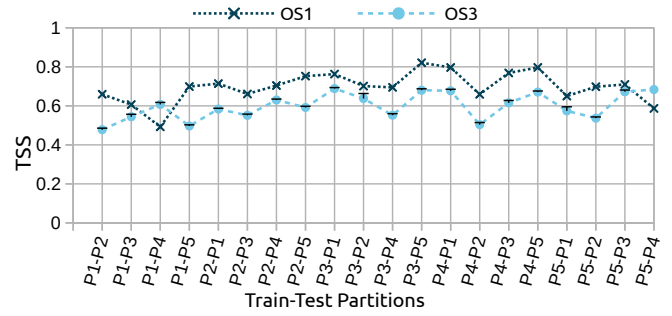


Fig. 8. [Experiment F] Average TSS of SVM performance impacted by two different oversampling methods from Table I: Oversampling 1 (OS1), where the climatology of sub-classes are preserved, and Oversampling 3 (OS3), where the sub-classes are forced to reach a 1 : 1 balance ratio by considering C-class to be the base. In all cases, global normalization is used.

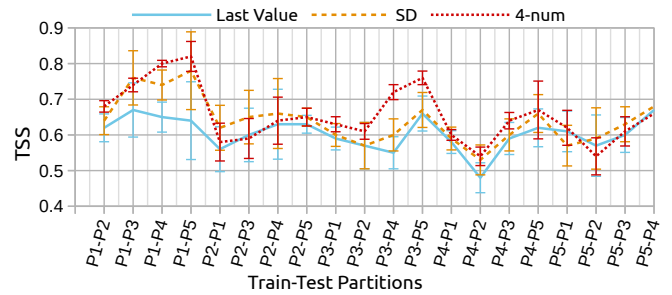


Fig. 9. [Experiment G] Average TSS of SVM performance on 3 different feature spaces: 1) *last value*, 2) *standard deviation*, and 3) *four-number summary* of time series, namely *median*, *standard deviation*, *skewness*, and *kurtosis*. Undersampling 2 (i.e., US2 from Table I) is used to remedy the class-imbalance issue. In all cases, global normalization is used.

Experiment G: SVM with Other Statistical Features.

During training in this experiment, we apply Undersampling 2 from Table I. We train and test SVM on all partition pairs using (i) *last value*, (ii) *standard deviation*, and (iii) *median*, *standard deviation*, *skewness*, *kurtosis*. As illustrated in Fig. 9, *standard deviation* results in statistically better performance than *last value*, and using the four-number summary seems to outperform *standard deviation*. This is a very good indication that different characteristics of time series carry some important pieces of information that may significantly improve reliability of a forecast model.

We leave the investigation for determining the optimal number of statistical features or selection of which statistical features are most effective for improving flare forecast performance for future studies.

VI. SUMMARY, CONCLUSIONS, AND FUTURE WORK

We used SWAN-SF benchmark dataset as a case study to highlight some of the challenges in working with imbalanced datasets, which are very often overlooked by scientists of the domain. We also addressed an interesting characteristic of some datasets, that we called *temporal coherence*, inherited from the spatial and temporal dimensions of the data. Using several different experiments, we showcased some pitfalls and

overlooked consequences of disregarding those peculiarities, and we discussed the impact of different remedies in the context of flare forecast problem.

There are still many interesting cases left to be discussed that we plan to include in our future studies. In several experiments, for instance, we noticed that despite the improvement in models' performances in terms of TSS, other measures such as HSS showed a moderate deterioration in the models. A new measure that reflects both these skill scores appears to be necessary to avoid many misleading interpretations. In spite of many studies, it is still an unsolved problem. Another avenue for further investigation is to incorporate NOAA AR Numbers in the sampling phase so that the temporal coherence of data that confines normalization and hyper-parameter tuning tasks to only subsets of data, can be bridged over.

We hope that this work raises awareness not only to the scientists interested in flare forecast problem, but all interdisciplinary researchers who might be dealing with imbalanced and temporally coherent datasets, and enables them to spot and tackle similar problems in their respective areas.

ACKNOWLEDGMENT

This work was supported in part by two NASA Grant Awards [No. NNX11AM13A, and No. NNX15AF39G], and one NSF Grant Award [No. AC1443061]. The NSF Grant Award has been supported by funding from the Division of Advanced Cyberinfrastructure within the Directorate for Computer and Information Science and Engineering, the Division of Astronomical Sciences within the Directorate for Mathematical and Physical Sciences, and the Division of Atmospheric and Geospace Sciences within the Directorate for Geosciences. Also, we would like to mention that all images used in this work are courtesy of NASA/SDO and the AIA, EVE, and HMI science teams.

REFERENCES

- [1] A. Ahmadzadeh and K. Sinha, "A multivariate time series data toolkit," Nov. 2019. [Online]. Available: https://bitbucket.org/gsudmlab/mvtsdata_toolkit/
- [2] B. Aydin *et al.*, "Multivariate time series dataset for space weather data analytics (manuscript submitted for publication)," *Scientific Data*, 2019.
- [3] M. J. Aschwanden *et al.*, "25 years of self-organized criticality: solar and astrophysics," *Space Science Reviews*, vol. 198, no. 1-4, pp. 47–166, 2016. [Online]. Available: <https://link.springer.com/article/10.1007%2Fs11214-014-0054-6>
- [4] G. Barnes *et al.*, "A comparison of flare forecasting methods. i. results from the "all-clear" workshop," *The Astrophysical Journal*, vol. 829, no. 2, p. 89, Sep 2016. [Online]. Available: <https://iopscience.iop.org/article/10.3847/0004-637X/829/2/89>
- [5] K. Florios, I. Kontogiannis, S.-H. Park, J. A. Guerra, F. Benvenuto, D. S. Bloomfield, and M. K. Georgoulis, "Forecasting solar flares using magnetogram-based predictors and machine learning," *Solar Physics*, vol. 293, no. 2, p. 28, Jan 2018. [Online]. Available: <https://doi.org/10.1007/s11207-018-1250-4>
- [6] A. Benz, "Flare observations," *Living Reviews in Solar Physics*, vol. 5, 02 2008. [Online]. Available: doi.org/10.12942/lrsp-2008-1
- [7] L. Fletcher, B. R. Dennis, H. S. Hudson, S. Krucker, K. Phillips, A. Veronig, M. Battaglia, L. Bone, A. Caspi, Q. Chen, P. Gallagher, P. T. Grigis, H. Ji, W. Liu, R. O. Milligan, and M. Temmer, "An observational overview of solar flares," *Space Science Reviews*, vol. 159, no. 1, p. 19, Aug 2011. [Online]. Available: <https://doi.org/10.1007/s11214-010-9701-8>
- [8] S. Tsuneta, K. Ichimoto, Y. Katsukawa, S. Nagata, M. Otsubo, T. Shimizu, Y. Suematsu, M. Nakagiri, M. Noguchi, T. Tarbell, A. Title, R. Shine, W. Rosenberg, C. Hoffmann, B. Jurcevic, G. Kushner, M. Levay, B. Lites, D. Elmore, T. Matsushita, N. Kawaguchi, H. Saito, I. Mikami, L. D. Hill, and J. K. Owens, "The solar optical telescope for the hinode mission: An overview," *Solar Physics*, vol. 249, no. 2, pp. 167–196, Jun 2008. [Online]. Available: <https://doi.org/10.1007/s11207-008-9174-z>
- [9] P. H. Scherrer, J. Schou, R. I. Bush, A. G. Kosovichev, R. S. Bogart, J. T. Hoeksema, Y. Liu, T. L. Duvall, J. Zhao, A. M. Title, C. J. Schrijver, T. D. Tarbell, and S. Tomczyk, "The helioseismic and magnetic imager (hmi) investigation for the solar dynamics observatory (sdo)," *Solar Physics*, vol. 275, no. 1, pp. 207–227, Jan 2012. [Online]. Available: <https://doi.org/10.1007/s11207-011-9834-2>
- [10] P. H. Scherrer, R. S. Bogart, R. I. Bush, J. T. Hoeksema, A. G. Kosovichev, J. Schou, W. Rosenberg, L. Springer, T. D. Tarbell, A. Title, C. J. Wolfson, and I. Zayer, *The Solar Oscillations Investigation — Michelson Doppler Imager*. Dordrecht: Springer Netherlands, 1995, pp. 129–188. [Online]. Available: https://doi.org/10.1007/978-94-009-0191-9_5
- [11] A. Tritschler, T. R. Rimmele, S. Berukoff, R. Casini, J. R. Kuhn, H. Lin, M. P. Rast, J. P. McMullin, W. Schmidt, F. Woger, and D. Team, "Daniel k. inouye solar telescope: High-resolution observing of the dynamic sun," *Astronomische Nachrichten*, vol. 337, no. 10, pp. 1064–1069, 11 2016. [Online]. Available: <https://onlinelibrary.wiley.com/doi/abs/10.1002/asna.201612434>
- [12] N. Fox, "Parker Solar Probe: A NASA Mission to Touch the Sun," in *EGU General Assembly Conference Abstracts*, vol. 20, Apr 2018, p. 10345.
- [13] M. G. Bobra and S. Couvidat, "Solar flare prediction using sdo/hmi vector magnetic field data with a machine-learning algorithm," *The Astrophysical Journal*, vol. 798, no. 2, p. 135, 2015. [Online]. Available: doi.org/10.1088/0004-637X/798/2/135
- [14] D. S. Bloomfield *et al.*, "Toward reliable benchmarking of solar flare forecasting methods," *The Astrophysical Journal*, vol. 747, no. 2, p. L41, Feb 2012. [Online]. Available: <https://iopscience.iop.org/article/10.1088/2041-8205/747/2/L41>
- [15] F. Benvenuto, M. Piana, C. Campi, and A. M. Massone, "A hybrid supervised/unsupervised machine learning approach to solar flare prediction," *The Astrophysical Journal*, vol. 853, no. 1, p. 90, Jan 2018. [Online]. Available: <https://doi.org/10.3847/1538-4357/aa23c>
- [16] F. Woodcock, "The evaluation of yes/no forecasts for scientific and administrative purposes," *Monthly Weather Review*, vol. 104, no. 10, pp. 1209–1214, 1976. [Online]. Available: [https://doi.org/10.1175/1520-0493\(1976\)104<1209:TEOYFF>2.0.CO;2](https://doi.org/10.1175/1520-0493(1976)104<1209:TEOYFF>2.0.CO;2)
- [17] G. Barnes and K. Leka, "Evaluating the performance of solar flare forecasting methods," *The Astrophysical Journal Letters*, vol. 688, no. 2, p. L107, 2008. [Online]. Available: <https://doi.org/10.1086/595550>
- [18] J. Mason and J. Hoeksema, "Testing automated solar flare forecasting with 13 years of michelson doppler imager magnetograms," *The Astrophysical Journal*, vol. 723, no. 1, p. 634, 2010. [Online]. Available: <https://doi.org/10.1088/0004-637X/723/1/634>
- [19] A. Ben-Hur and J. Weston, "A user's guide to support vector machines," *Methods in molecular biology (Clifton, N.J.)*, vol. 609, pp. 223–39, 01 2010. [Online]. Available: [doi.org/https://doi.org/10.1007/978-1-60327-241-4_13](https://doi.org/10.1007/978-1-60327-241-4_13)