

Improving earthquake monitoring for gravitational-waves detectors with historical seismic data

SKY SOLTERO¹

¹*Mentor: Michael Coughlin*

ABSTRACT

A remarkable level of isolation from the ground is required for Advanced gravitational-wave detectors such as the Laser Interferometer Gravitational-Wave Observatory (LIGO) to function at peak performance. These ground based detectors are susceptible to high magnitude teleseismic events such as earthquakes, which can disrupt proper functioning, operation and significantly reduce their duty cycle. As a result, data is lost and it can take several hours for a detector to stabilize and return to the proper state for scientific observations. With advanced warning of impending tremors, the impact can be suppressed in the isolation system and the down time can be reduced at the expense of increased instrumental noise. An earthquake early- warning system has been developed relying on near real-time earthquake alerts provided by the U.S. Geological Survey (USGS) and the National Oceanic and Atmospheric Administration (NOAA). The alerts can be used to estimate arrival times and ground velocities at the gravitational-wave detectors. By using machine learning algorithms, a prediction model and control strategy has been developed to reduce LIGO downtime by 30%. This paper presents further improvements under consideration to better develop that prediction model and decrease interruptions during LIGO operation.

I. INTRODUCTION

The two detectors that compose the Laser Interferometer Gravitational-Wave Observatory (LIGO) along with Virgo, and GEO600 detectors form a global network of gravitational wave interferometers. Keeping the detectors in operating mode requires an exceptional level of isolation from the ground so that the cavities can be held in optical resonance and be capable of observing displacements in space-time of less than one thousandth of the diameter of a proton. Environmental disturbances such as earthquakes can disrupt operating mode, destabilize detectors and cause the detectors to fall out of lock despite seismic isolation systems already in place to minimize interfering effects. When the detectors have fallen out of lock, where the control system cannot maintain optics at their stabilized positions, it can take many hours to return to the locked state and normal operation. During the observation run (referred as O1), from January 18, 2015 to January 12, 2016 operation was disrupted 62 times at LIGO Hanford and 83 times at LIGO Livingston due to earthquakes. Previous studies have shown that by using an early-warning earthquake system, relying on alerts provided by the U.S. Geological Survey (USGS) and the National Oceanic and Atmospheric Administration (NOAA), arrival times and ground velocities could be predicted which have a direct correlation with the operation status of the interferometers (Coughlin et al. 2017). The higher the incoming

seismic velocities the more unstable the interferometer. A strategy intended to maintain lock and suppress these seismic disturbances early in the isolation system, at the expense of sensitivity and increased noise, would notably increase the interferometers' duty cycle (Biscans et al. 2018). Consequently, an earthquake early warning application named Seismon has been created to process real-time alerts from the USGS containing specific characteristic information about the earthquakes to provide estimated arrival times of the seismic phases and seismic amplitudes of the surface waves at the detector sites. By implementing detector control configurations, it is predicted that 40 to 100 earthquake operation interruptions could be prevented in a 6-month period.

II. OBJECTIVES

We aim to improve the algorithms of Seismon and as a result reduce LIGO downtime and increase the time the detectors are in observing mode. The alerts received from USGS contain information on time, location, depth, and magnitude of a specific earthquake which is then used to predict ground velocities, arrival time and amplitude of the various seismic phases at the detector sites. Seismon initially relies on earthquake notifications from a worldwide network of seismometers. P-waves (primary) traveling twice as fast as S-waves (secondary) reach the seismic stations first, thus providing the initial earthquake character estimations. As more and more data is acquired solutions to

the hypocenter and magnitude of the earthquake are estimated and the solutions are sent to USGS’s Product Distribution Layer (PDL). This ensures Seismon receives the most pertinent notifications. From there the notifications are processed to predict the seismic wave arrival time and the amplitude of the ground motion at the detectors. Past earthquake records and the seismic data at the detectors are also examined to predict how the ground motion will affect the observatories. The predicted amplitude and past earthquake data are compared, with the difference being minimized by adaptive simulated annealing algorithms to obtain solutions close to the global minima. Lastly, the predictions are used to create warnings delivered to the detectors containing the amplitude prediction, lockloss probability and the anticipated earthquake arrival time at the observatories. Seismon performance can be evaluated by recording and analyzing the notification duration, accuracy of predicted ground-motion amplitude, time-of-arrival predictions and the detector lockloss predictions. Current evaluations with the LIGO Observing Run 1 from September 2015 to January 2016, show about 90% of seismic events are within a factor of 5 of the predicted ground velocity and within 3s of the final predicted arrival time (Coughlin et al. 2017). Examining the times lockloss occurred, it can be said that the detectors generally fall out of lock at ground velocities greater than $5 \mu\text{m/s}$ but at lower velocities the data is more complex. Therefore, incorporating more ways of determining better lockloss predictions are of interest and would demonstrate success in this project. We purpose to improve the Seismon algorithm by incorporating more machine learning methods, broadening ground motion parameters and collecting more accurate data to enrich the prediction models.

III. APPROACH

We intend on advancing the Seismon application by improving predictions and acquiring more data of various parameters of incoming teleseismic events. We will test if the arrival time predictions can be improved by machine learning algorithms. To enhance ground velocity predictions, we will explore broadening our data resources and determine if we can acquire more data from hundreds of other seismic stations around the United States and the world.

IV. METHODS

A. INVESTIGATING P-WAVES

To better understand the effects of earthquake magnitude and global location on arriving earthquake surface velocities at the detectors, P-wave properties and ve-

locities were initially investigated. Multiple plots using historical data from Livingston and Hanford gathered from observation run 1 and 2 have been made. In Figure 1, earthquake velocities are determined by dividing the distance from earthquake origin to detector by the difference of P-surface wave prediction times and earthquake times. These earthquake velocity magnitudes are then plotted at their origin in regards to latitude and longitude. These plots show the higher velocity P-wave are distinguished clearly in the historical data.

Typically P-wave velocities are twice the speed of surface waves yet these plots show velocities reaching up to 16,000 m/s, which is higher than expected for simple P-waves. Higher velocity speeds are often due to a certain azimuth degree and dept of the earthquake origin which has an effect on the internal reflections the wave experiences while traveling within the Earth. To explore the contributions of reflections internal to the Earth, in Figure 2, the effective velocities of the paths that include reflections (left) and those that do not (right) are shown. These velocities are shown on a grid of depth and azimuth degrees. While the plot on the right is in line with expectations for P-wave velocities, the plot on the left shows the contributions from reflections, leading to much higher effective velocities (and therefore faster arrivals). This result shows that the first arrivals of the P-waves shown in Figure 1 derive from P-waves reflecting in the Earth and that the historical data collected at the LIGO interferometers includes P-wave data inclusive of these reflections.

The surface velocities derived from the same historical data at the Hanford and Livingston sites can now be explored. For easy comparison the same world maps were made with surface wave velocity data in Figure 3 3 which show the effective earthquake velocity, measured as the distance of the earthquake divided by the difference of the peak ground velocity time and earthquake time. It shows a range of velocities from 2,000 to 5,000 m/s which is appropriate for surface wave velocities dominating the time-series, as expected. We include only the historical data with peak ground velocities greater than $1 \mu\text{s}$.

A further analysis in the form of a histogram shows the frequency of earthquakes at certain velocities. Figure 4 shows a histogram corresponding to the data used for Figure 3. These plots show that the majority of earthquakes have effective velocities between 2,000 and 4,000 m/s, as expected for surface waves. However, this graph shows that there are high speed outliers likely due to body wave contributions such as a P-waves or contamination from other earthquakes. This contamination of earthquake waves different from simple surface waves yet being categorized as surface waves can interfere with

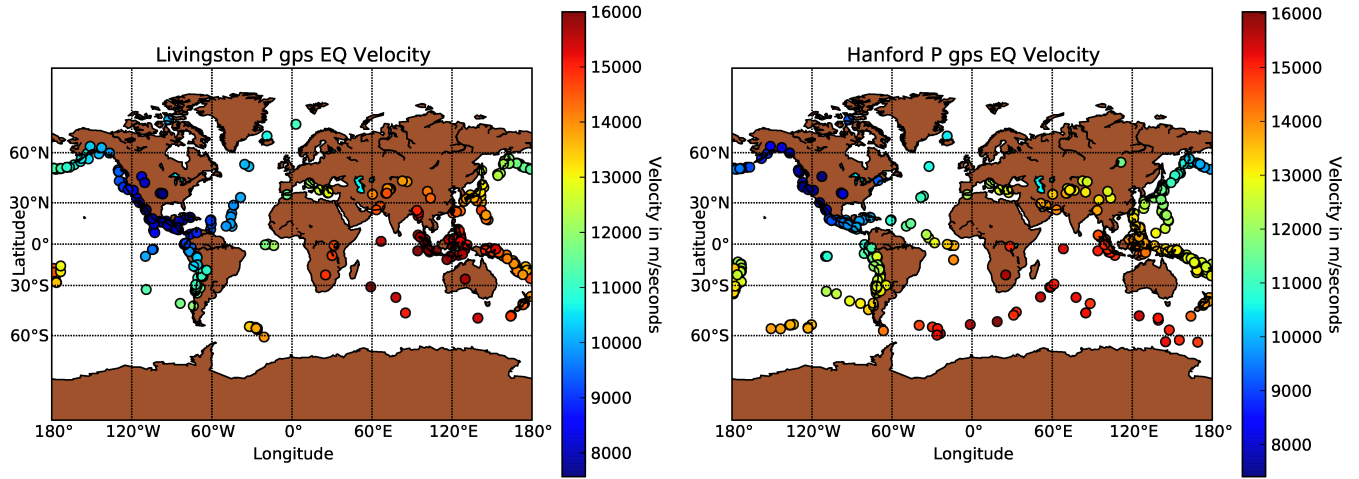


Figure 1. Magnitude of earthquake (EQ) velocities based on data using P-wave arrival times plotted at corresponding latitude and longitude points. Data with peak ground velocities less than $1\mu/s$ have been omitted. Displayed on the left is the plot for Livingston data and on the right is the plot for Hanford data.

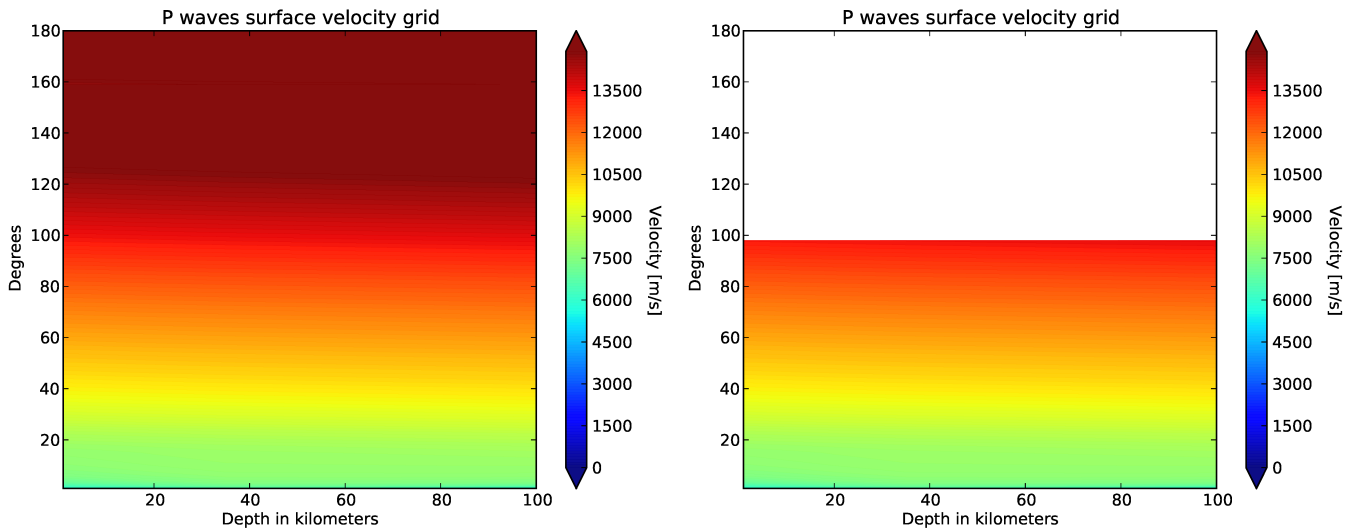


Figure 2. Magnitude of Earthquake velocities based on data using P-wave arrival times plotted in accord to degrees and dept of earthquake origin. Displayed on the left plot is the velocity without taking into account reflections. On the right plot reflections are taken into account and more expected P-wave velocities are shown.

the accuracy of surface wave predictions used in machine learning algorithms.

B. NEURAL NETWORK PREDICTIONS

Different experimentation in machine learning also led to some practical results. Using a basic Keras neural network, with a loss set to mean square error, historical data from Livingston, Hanford and Virgo seismometers were used to create and test prediction models for ground velocity and surface wave arrival time estimates. A similar neural network with the loss set to categorical crossentropy was used to predict detector status. The neural network took 6 input parameters that consisted of earthquake magnitude, longitude, latitude, distance from detector, dept and azimuth degree. With only these 6 parameters initially available the goal was to train a neural network model that could accurately make the above predictions. When training the neural network a randomized portion of data is set as the training set and another portion as the testing set. Experimentation with the amount of hidden layers and various hyper-parameters such as batch size, learning rate, dropout rate, activation etc. was explored and the best result models were saved. The predictions using the historical data and in particular for the arrival time predictions, all fell within a factor of 3.5 of the actual value. This factor of 3.5 is determined by taking the difference of the predicted and known actual ground velocities and then dividing by the actual.

To get better neural network results, broadening data resources was necessary. We decided to utilize earthquake data collected at Incorporated Research Institutions for Seismology (IRIS), which included over 730,000 Earthquakes compared to the small amount of 2,000-3,000 earthquakes for the historical datasets. The same six input parameters were used to predict ground velocities and surface wave arrival times. This greatly improved the accuracy of predictions but we noticed the IRIS data included typical surface wave velocities and higher velocities intermixed. To check the type of earthquake data within the IRIS dataset being fed into the model is accurate and useful in itself, a plot of distance versus time was produced in Figure 5. The three labeled lines plotted show where known Surface rayleigh waves at different velocities would be represented in this graph. A corresponding graph visualizes these points in the form of density. These plots show us a high amount of velocities characteristic of P-wave velocities. Therefore a filter was placed on the IRIS data to only include data with velocities less than 6,000 m/s and then train the neural network model. This further improved the accuracy of surface wave predictions and in Figure 6,

the actual versus predicted arrival times were plotted in a density plot on the left. This plot shows the majority of earthquakes having arrival times between 2,000 and 6,000 seconds, concentrated along the trend line. We expect to see the predicted and actual arrival times plotted following a slope of 1. The more fitting to the slope line shown, the more accurate the predictions. On the right the counts versus relative error were plotted in a histogram. The histogram shows all the arrival time predictions within a factor of 1 with the majority within a factor of .4. Using the IRIS data proved to improve the neural network accuracy for ground velocity and arrival time predictions greatly. To predict lockloss status, a combination of Hanford and Livingston data was used to train a neural network model based on a loss of categorical crossentropy. The same six input parameters were used but the output was a prediction label of 0, 1, 2. The label 0 corresponds to detector not locked, 1 locked but can't take data and 2 Locked but loss lock.

V. PERFORMANCE

Now that arrival times could be accurately predicted under a factor of 1 using IRIS data we wanted to test that particular model on the other historical data sets of Livingston, Hanford and Virgo using the same six input parameters. Testing all of Hanford data, including velocities greater than 6,000 m/s the model was used to predict arrival times and did so within a factor of 3.5 as can be seen in Figure 7 on the left. However, if the Hanford data was filtered beforehand to exclude data with earthquake velocities greater than 6,000 m/s the predictions were slightly more accurate and is shown in the right plot. Similar results were produced for ground velocities. The lockloss model had the most inaccurate predictions. A randomized portion of the combined Hanford and Livingston data was set at the training set and another as the testing set. The results were then plotted in a confusion matrix in Figure 8. The quantity of true label versus the label it was predicted as is displayed on the left. Since the amount of different labels are not proportional the graph was transformed into a normalized matrix on the right that shows the percentage of that label predicted. It is shown that only 53% of the 0 label, 50% of the 1 label and 17% of the 2 label were predicted accurately.

VI. CONCLUSION

In this paper we have discussed the effects of earthquakes on gravitational-wave observatories and a earthquake early warning application named Seismon in place already to combat the effects of earthquakes. We further discuss improving a portion of Seismon's algorithm

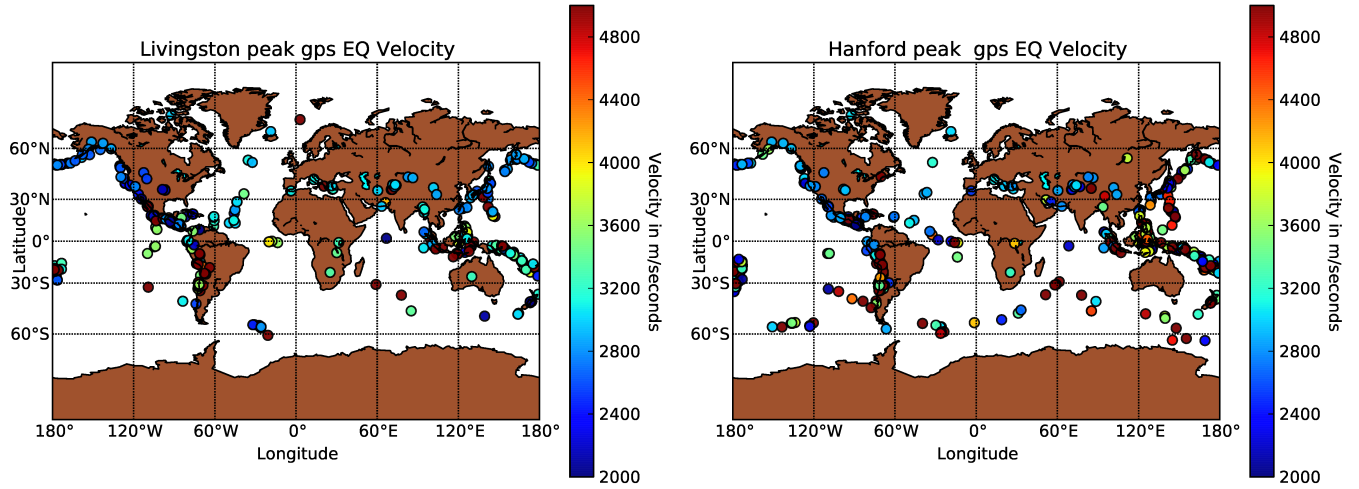


Figure 3. Magnitude of earthquake (EQ) velocities based on data using peak ground velocity gps time plotted at corresponding latitude and longitude points. Data with peak ground velocities under $1\text{e-}6$ have been omitted. Displayed on the left is the plot for Livingston data and on the right is the plot for Hanford data.

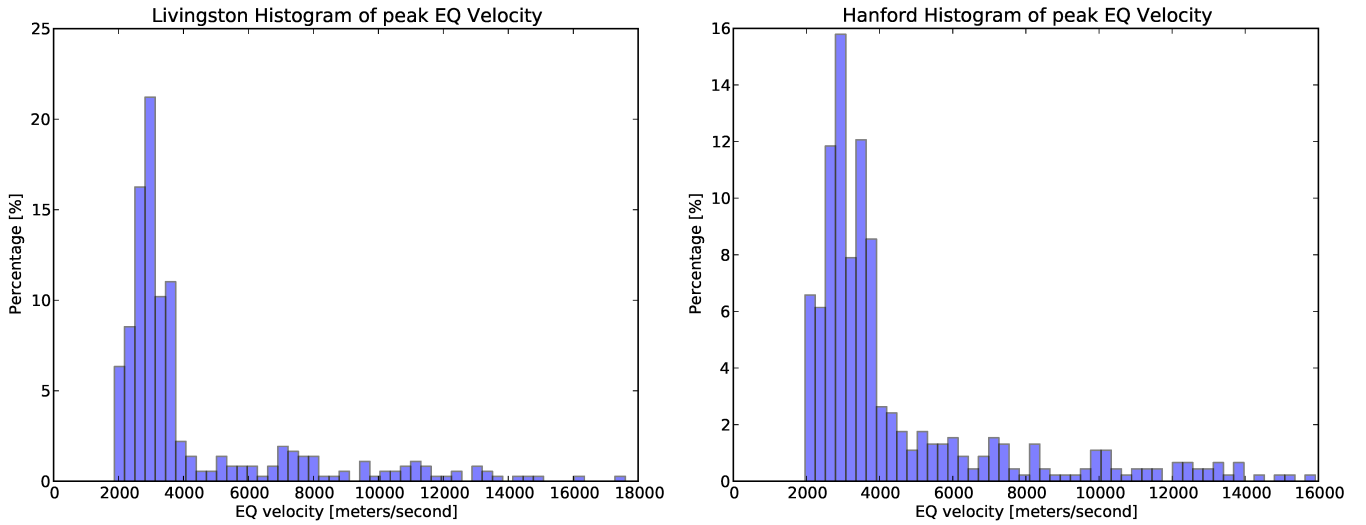


Figure 4. Percentage of different earthquake (EQ) velocities based on data using peak ground velocity gps time divided by the distance from the detectors. In association with the above Figure 3 plots. Data with peak ground velocities less than $1\text{ }\mu\text{/s}$ have been omitted. Displayed on the left is the plot for Livingston data and on the right is the plot for Hanford data.

in predicting earthquake ground velocities, surface wave arrival times and detector lockloss status. We have shown the neural network models for ground velocity and arrival times can be trained to make predictions accurately within a factor of 1 but do not predict as accurately on other historical datasets. It is shown we can predict lockloss status as well, but increased amount of training data would be beneficial. In the future we

hope neural network models can be implemented into the seismon pipeline to increase the time the various gravitational-wave detectors are in observation mode and the amount of gravitational-waves detected. Further efforts will be needed to determine what type of data range is best and useful to train the models and then be tested on. Further investigation into the lockloss status model and its correlation to input parameters is also recommended.

REFERENCES

- Biscans, S., Warner, J., Mittleman, R., et al. 2018, Classical and Quantum Gravity
- Coughlin, M., Earle, P., Harms, J., et al. 2017, Classical and Quantum Gravity, 34, 044004

The authors are grateful for the support of the US National Science Foundation's Research Experience for Undergraduates Program, award #1757303.

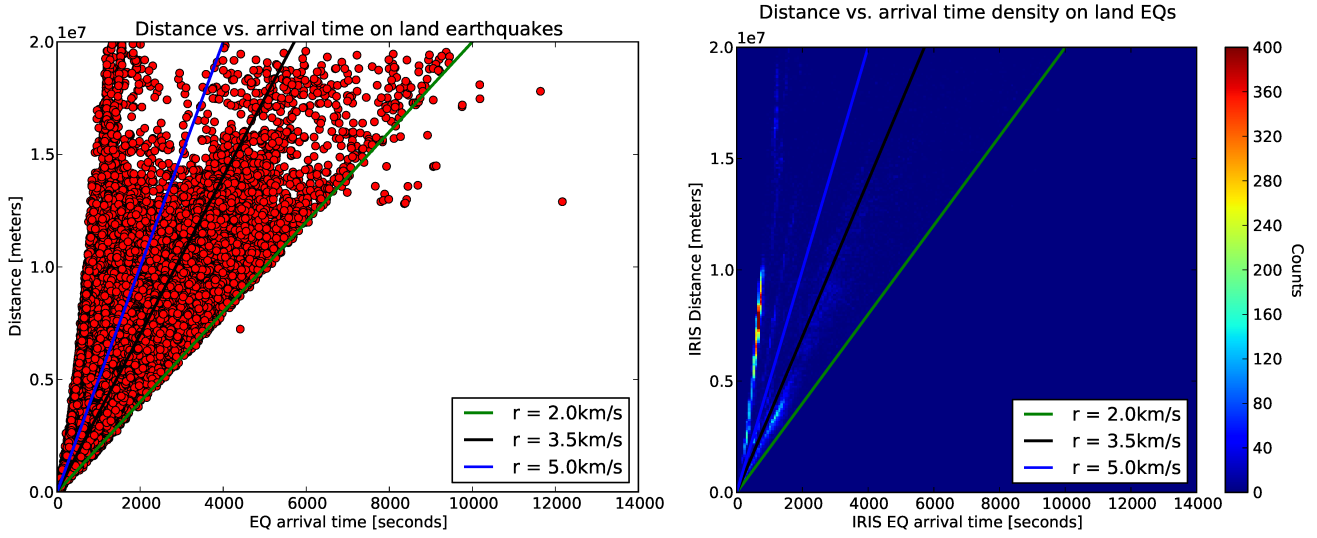


Figure 5. The distance versus arrival times of the input data used to train the neural network model and produce Figure 6 graphs. The solid slope lines represent different rayleigh wave speeds commonly observed. A density plot of the same graph for easier visualization of data is shown on the right. Data with peak ground velocities less than $1\mu\text{s}$ have been omitted.

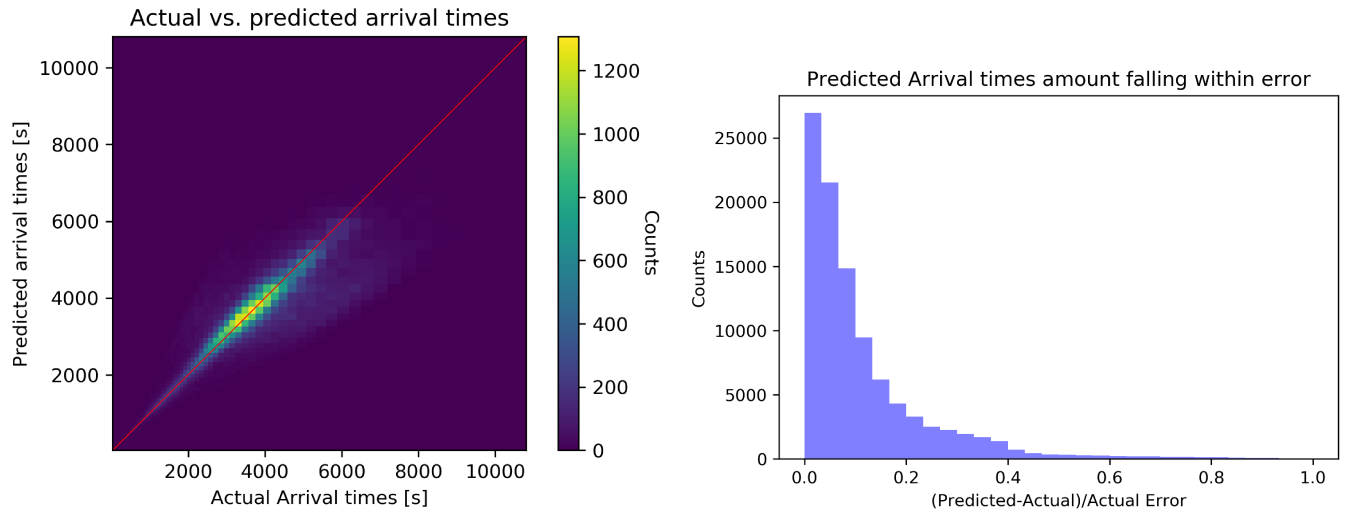


Figure 6. Surface wave arrival time predictions versus actual arrival times density plot (left). The closer the points displayed are to the trend line of slope 1, the more accurate the prediction. Counts of earthquake arrival times versus the relative error (right). Data with peak ground velocities less than $1\mu\text{s}$ have been omitted.

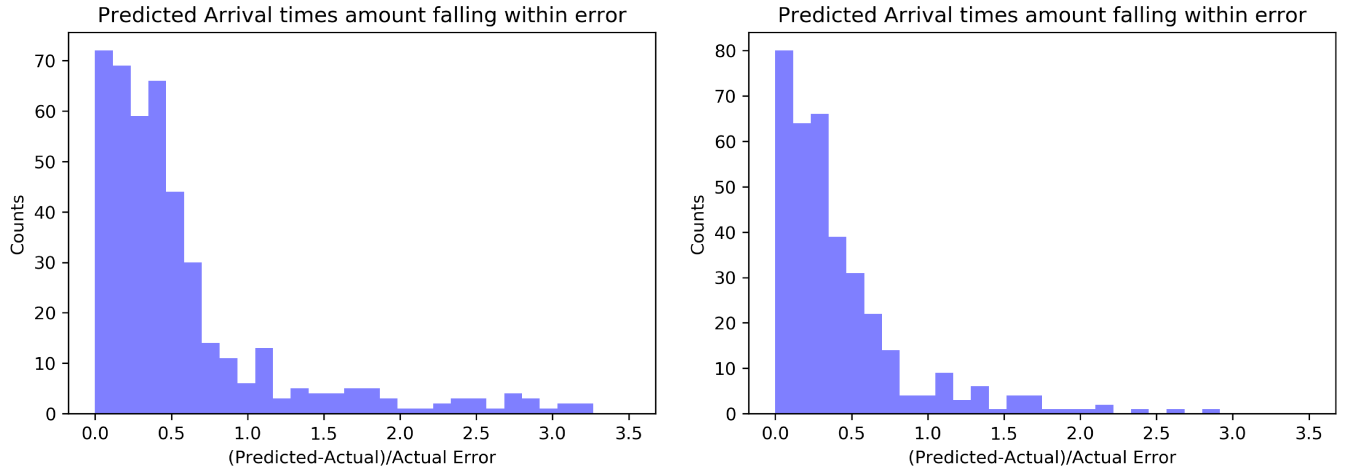


Figure 7. Counts of earthquake arrival times versus the relative error using all Hanford data (left). Counts of earthquake arrival times versus the relative error using Hanford data omitting velocities greater than 6,000 m/s (right). Data with peak ground velocities less than $1 \mu/s$ have been omitted.

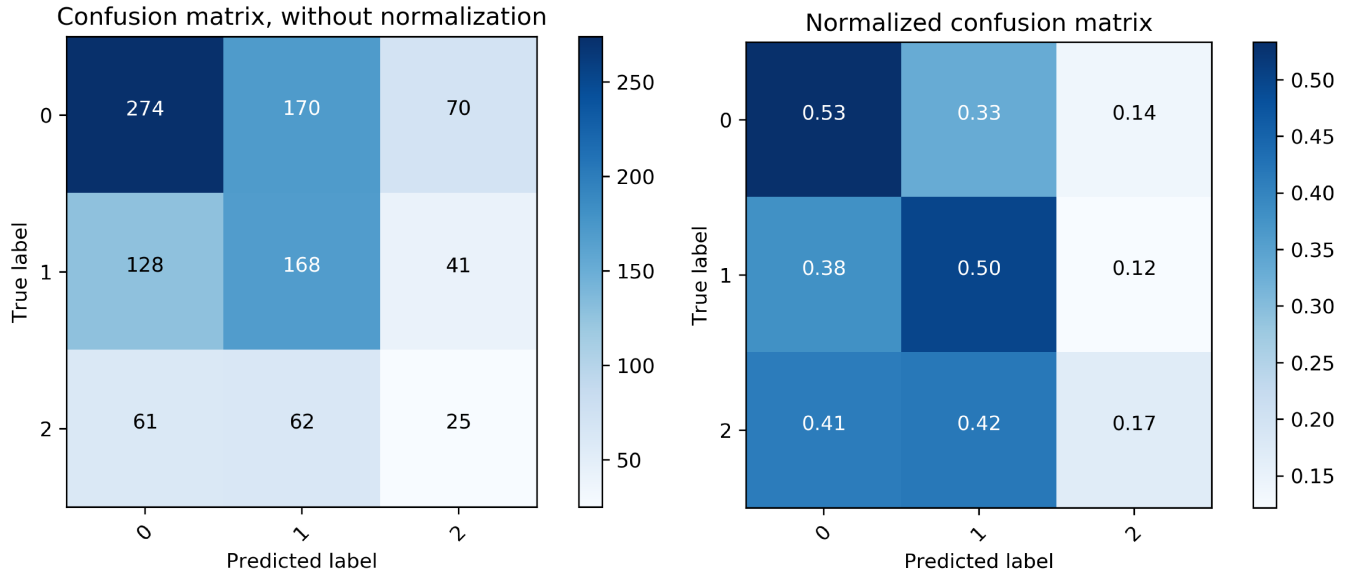


Figure 8. A matrix showing counts of the actual lockloss status label versus the predicted label (left). The right shows the same plot normalized to show the amount in percentage of a certain label and its prediction. The label 0 corresponds to detector not locked, 1 locked but can't take data and 2 locked but loss lock. Data with peak ground velocities less than $1 \mu/s$ have been omitted.