



Application of Machine Learning Techniques To Predict Wildfire Risk Using Satellite Imagery

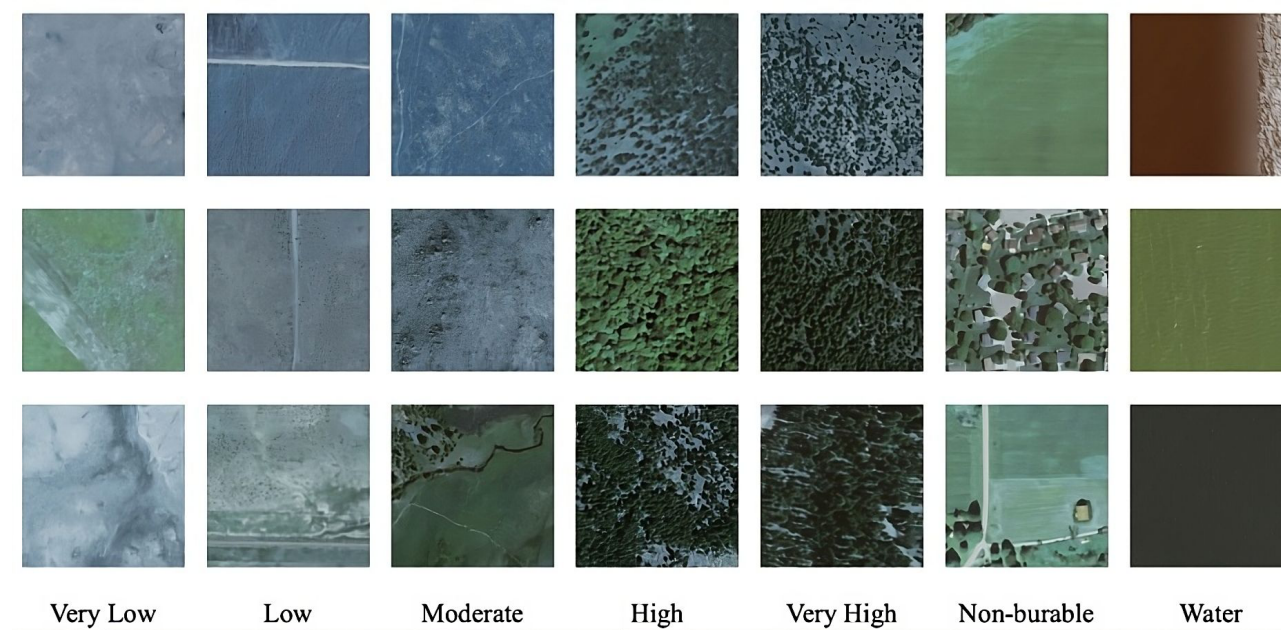
Anisha Palaparthi¹, Nachiketh Karthik¹

anishapv@stanford.edu, karthi24@stanford.edu

Stanford
Computer Science

Project Overview

- Increasing wildfire frequency, especially in California, has driven the need for improved risk assessment.
- Current machine learning models rely heavily on complex data and deep learning, which require advanced tools and significant computational resources.
- We aim to address this by combining simpler unsupervised and supervised learning models together in a pipeline to achieve similar performance as deep learning models.
- Applying our proposed method, we achieve non-trivial performance compared to the baselines with 37% accuracy

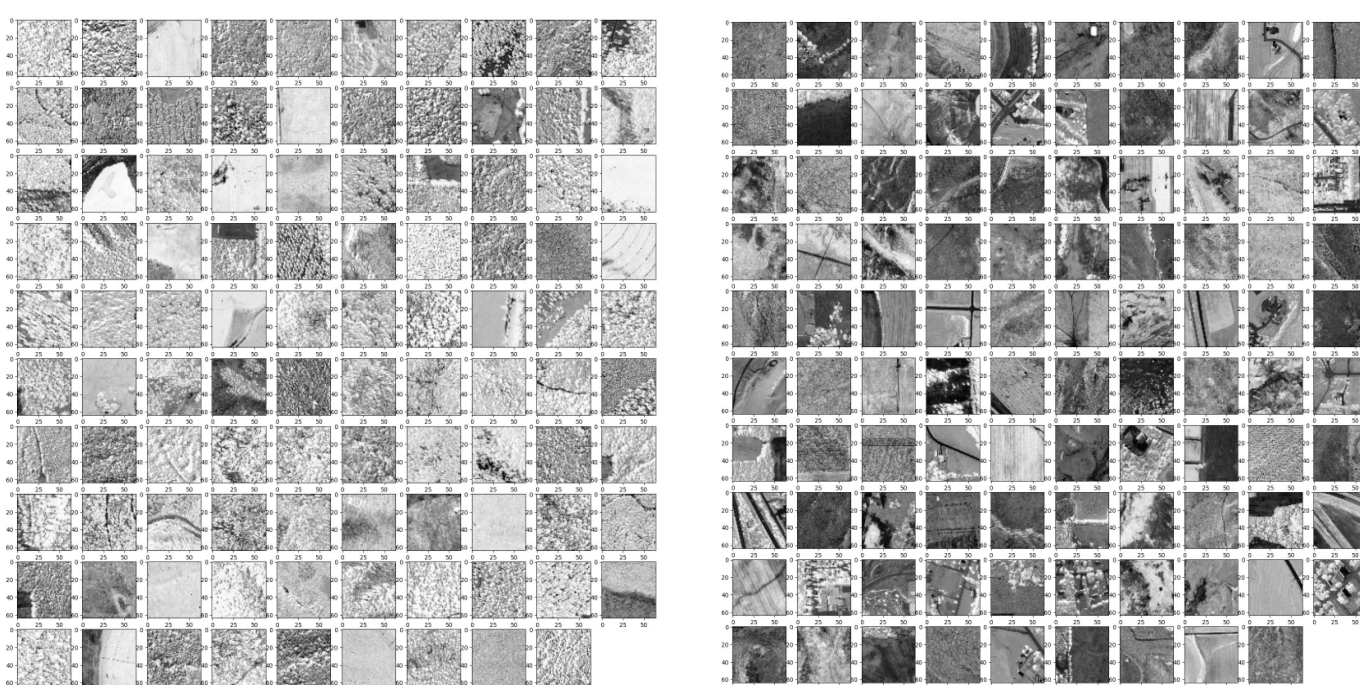


Data and Features

- Data obtained from FireRisk dataset [1] and contains RGB satellite images of size 320x320x3.
- Labelled with ground truth of Wildfire Risk level: "Very Low"-"Very High", "Non-burnable", "Water" (7 classes).
- Converted to B&W 64x64 image, reduced to 400 most salient features using PCA.
- Use K-Means to cluster similar images - e.g. k=3 includes grass/trees, k=4 includes roads

Cluster #3 Visualization

Cluster #4 Visualization



Models

SVM with RBF kernel

$$\min_{\mathbf{w}, b} \frac{1}{2} \|\mathbf{w}\|^2 + C \sum_{i=1}^N \max(0, 1 - y_i(\mathbf{w}^T \mathbf{x}_i + b))$$

- C [Regularization] = 1
- Gamma [Kernel Coefficient] = 'Scale'

$$\gamma = \frac{1}{n_{\text{features}} \cdot \text{Var}(X)}$$

K-Means Algorithm

$$J = \sum_{i=1}^k \sum_{x \in C_i} \|x - \mu_i\|^2$$

- K=5 clusters
- Update step

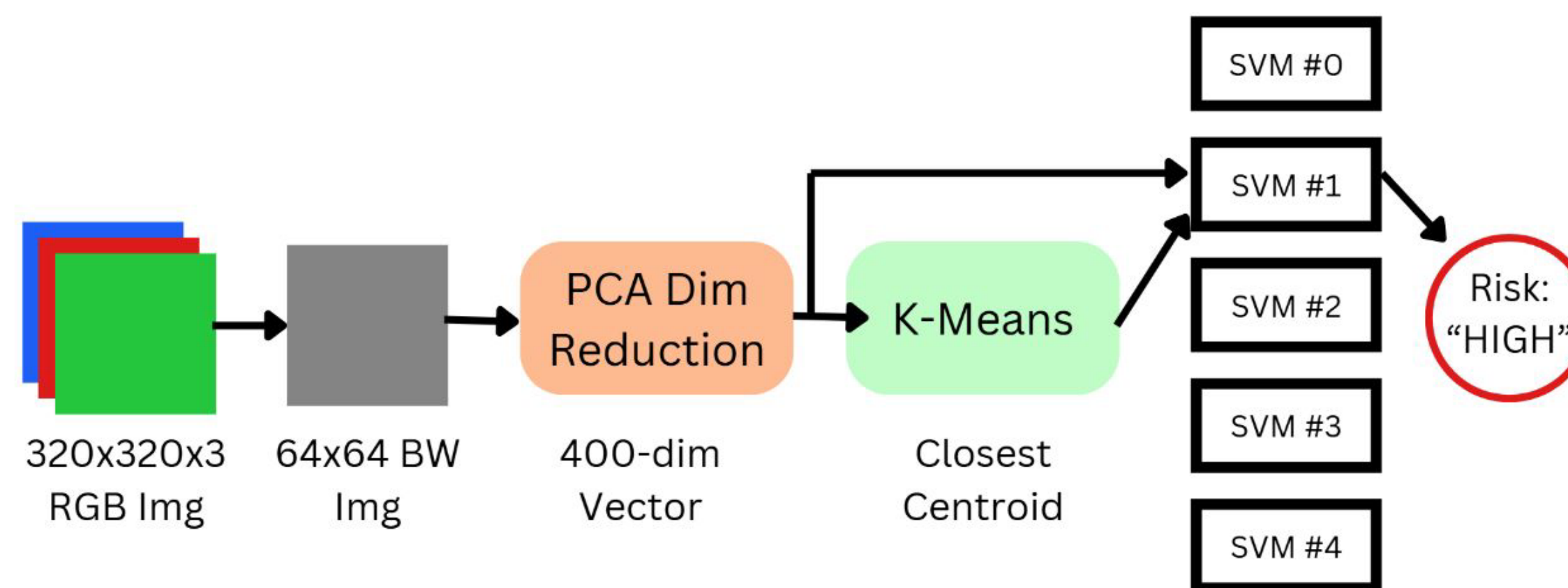
$$\mu_i = \frac{1}{|C_i|} \sum_{x \in C_i} x$$

Principal Component Analysis

$$x_{\text{pca}}^{(i)} = \begin{bmatrix} u_1^T x^{(i)} \\ u_2^T x^{(i)} \\ \vdots \\ u_n^T x^{(i)} \end{bmatrix}$$

- d= 4096, n= 400
 - ~86% Variance Explained
- $$x^{(i)} \in \mathbb{R}^d, x_{\text{pca}}^{(i)} \in \mathbb{R}^n$$

Experiments

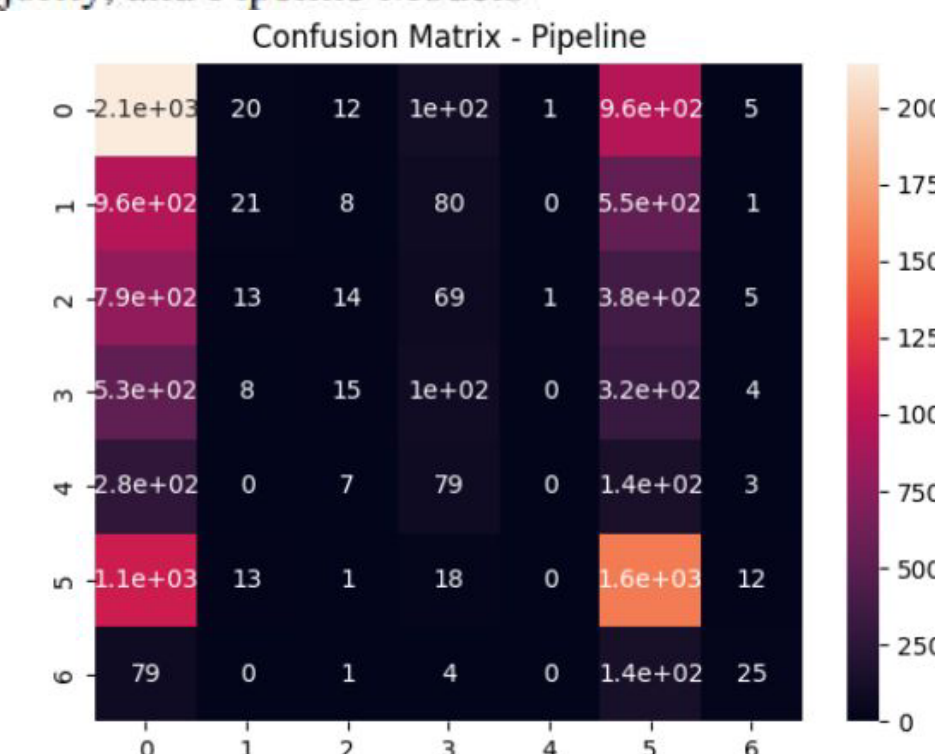


Class	Baseline			Majority			Pipeline		
	Precision	Recall	F1-Score	Precision	Recall	F1-Score	Precision	Recall	F1-Score
Very_Low	0.36	0.27	0.31	0.31	1.00	0.47	0.37	0.66	0.47
Low	0.09	0.03	0.04	0.00	0.00	0.00	0.28	0.01	0.02
Moderate	0.14	0.02	0.03	0.00	0.00	0.00	0.24	0.01	0.02
High	0.00	0.00	0.00	0.00	0.00	0.00	0.22	0.10	0.14
Very_High	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00
Non-burnable	0.29	0.75	0.42	0.00	0.00	0.00	0.38	0.58	0.46
Water	0.00	0.00	0.00	0.00	0.00	0.00	0.45	0.10	0.17
Accuracy	0.29			0.31			0.37		

Table 6: Comparison of Classification Reports for Baseline, Majority, and Pipeline Models

Training and Test Set Info

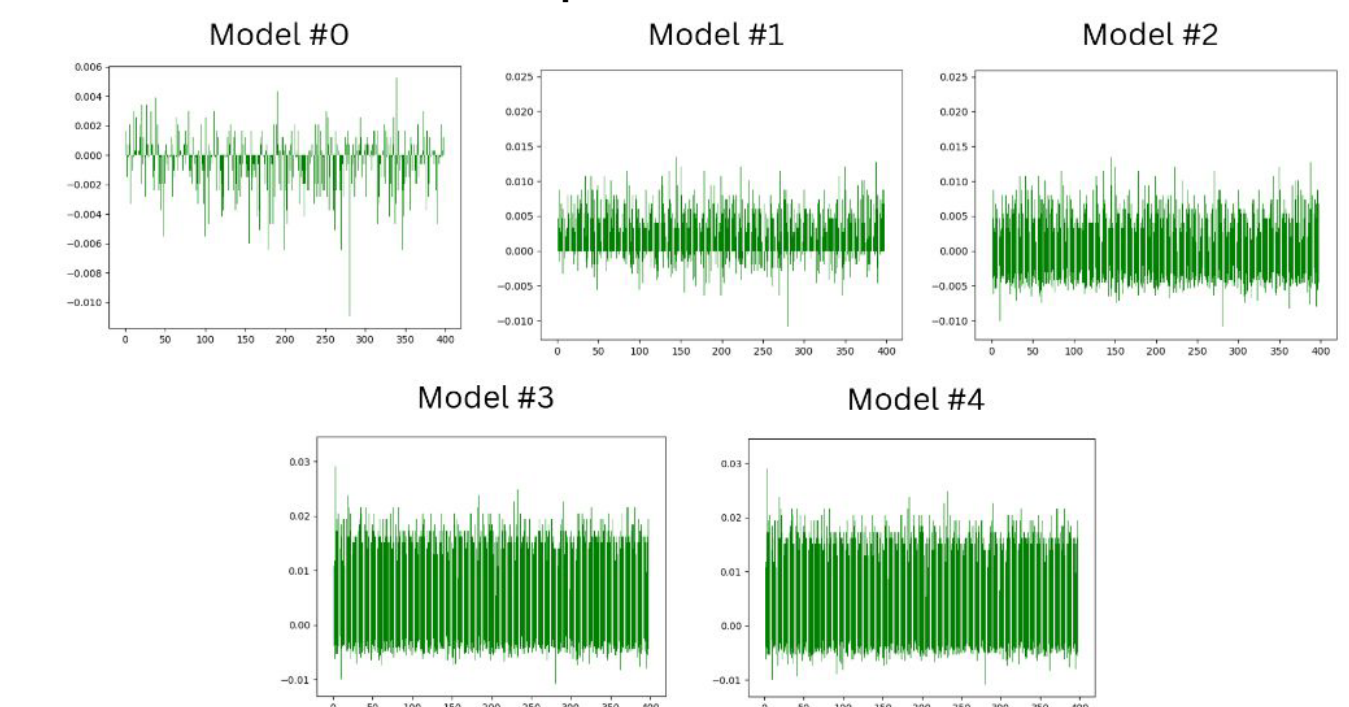
- Total Training Set Images: 70,331
- Total Validation Set Images: 21,541
- All three models are trained with 50% of the dataset (approx. 35K sample images) and evaluated on an unseen test set of that is 15% of the total dataset (approx. 10K images)



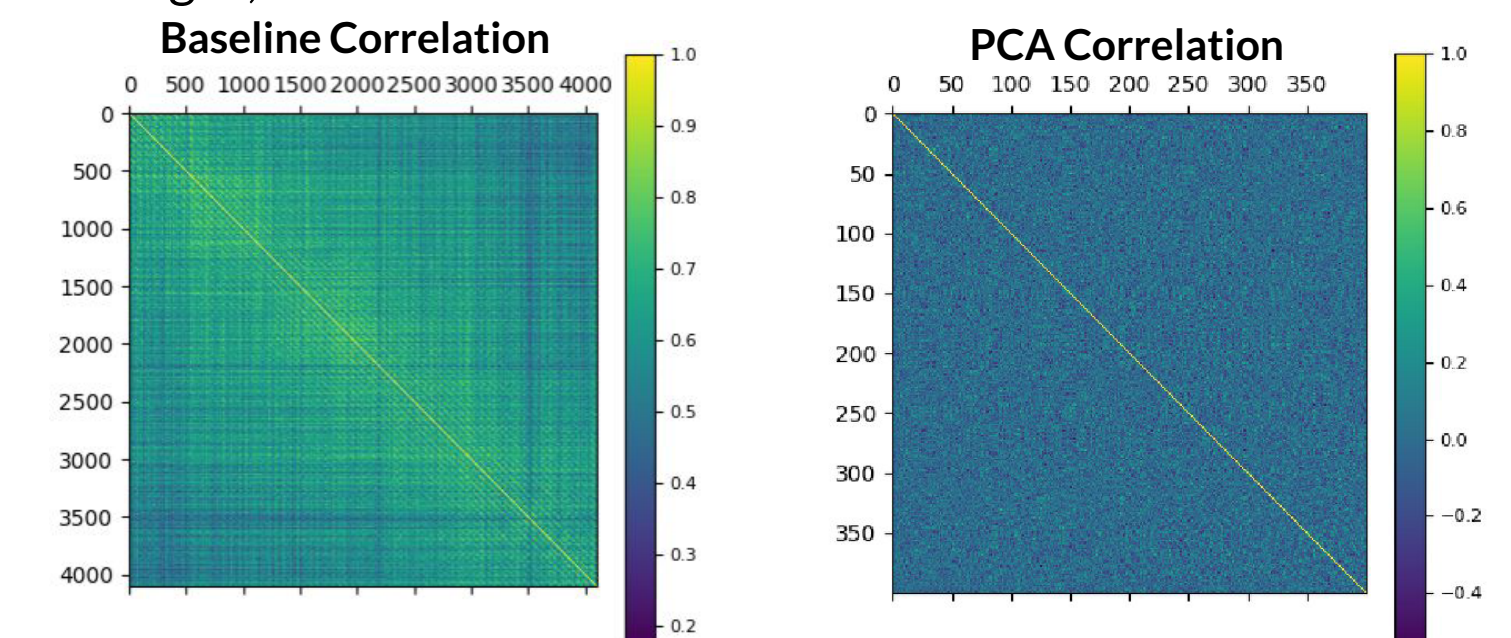
Discussion

- Achieves non-trivial improvement over baseline exps - 37% accuracy. Still struggles to classify all images (potentially due to input data loss in pipeline), particularly in the "Very High" risk class

Feature Importances Across K SVMs



- The importance of a given feature changes depending which model is evaluating that data, and is especially apparent in Model #0.
- Expected behavior since each model is trained on a different cluster of similar images, so each learns to focus on different features



- Many of the features (pixel values) are highly correlated with each other in the baseline which contributes, in part, to the low performance of the baseline model. On the other hand, PCA substantially reduces the feature correlation, enabling better performance of the SVM

Future Research

- We would like to explore what other machine learning models would benefit from the pipeline implemented here.
- We would also explore if the same principles of clustering and dimensionality reduction in our unsupervised step could potentially improve the performance of deep learning systems as well.

References

[1] Shuchang Shen, Sachith Seneviratne, Xinye Wanyan, and Michael Kirley. 2023. Firerisk: A remote sensing dataset for fire risk assessment with benchmarks using supervised and self-supervised learning.